

MASARYKOVA UNIVERZITA



Sborník příspěvků

9. letní škola aplikované informatiky

Editoři

Jiří Hřebíček

Jan Ministr

Tomáš Pitner

Bedřichov, 3.– 5. září 2012

Brno 2012

Slovo úvodem

9. letní škola aplikované informatiky navázala na předchozí letní školy aplikované (environmentální) informatiky v Bedřichově, které se zde konají od roku 2002 s výjimkou roku 2004, kdy se letní škola konala v Šubířově a roku 2005, kdy byla hlavní akcí Masarykovy univerzity (MU) v oblasti environmentální informatiky na 19. mezinárodní konferenci *Informatics for Environmental Protection - EnviroInfo 2005* s nosným tématem *Networking environmental information* a kterou hostila Masarykova univerzita ve dnech 7. až 9. září 2005 v brněnském hotelu Voroněž. V letech 2006 až 2011 se letní školy konaly opět v Bedřichově v pensionu U Alexů.

V letošním roce 9. letní škola aplikované informatiky se konala ve dnech 3. až 5. září 2012 v Bedřichově, kde proběhl workshop věnovaný prezentaci scénáře „*Anthropogenic Impact and Global Climate Change*“ řešeného Masarykovou universitou v rámci řešení projektu 7. rámcového programu Evropské unie č. 247893 „*TaToo - Tagging Tool based on a Semantic Discovery Framework*“. Tento projekt se zaměřuje na vývoj nástrojů Jednotného informačního prostoru v Evropě pro životní prostředí (Single Information Space in Europe for Environment - SISE) umožňující uživatelům snadno zjistit environmentální zdroje na webových stránkách (data, informační služby a modely, které mají různé informační uzly) a doplnit je o cenné informace v podobě sémantických anotací těchto zdrojů, což usnadní jejich budoucí použití a nalezení, a zahájí cyklus prospěšného obohacení environmentálních zdrojů. Navrhovaný rámec projektu *TaToo* je obecné povahy a umožní začlenění sémantiky, s přihlédnutím k různým návrhům doménových ontologií v environmentálních multidoménách a vícejazyčných souvislostech. Řešení projektu *TaToo* poskytne tři komplexní a rozsáhle ověřené scénáře, a proto se předpokládá, že jako hlavní cílová skupina uživatelů budou kvalifikovaní odborní uživatelé těchto scénářů z ústavů Institut biostatistiky a analýz (IBA) a Centrum pro výzkum toxických látek v prostředí (RECETOX) MU.

Všemi přednáškami 9. letní školy prolínala skutečnost, že aplikace moderních informačních a komunikačních technologií pro životní prostředí (potažmo environmentální informatiky) jak v České republice (ČR), tak i mezinárodně v Evropské unii (EU) a ve světě se zaměřuje na podporu eEnvironmentu, Jednotného informačního prostoru pro životního prostředí (SISE – Single Information Space in Environment for Europe) a Sdíleného informačního systému pro životní prostředí (SEIS – Shared Environmental Information System), které podporují naplňování nové politiky v budování informační společnosti EU a ČR, které přinesla „*Digital Agenda for Europe*“ v rámci vize nové Evropské komise „*eEUROPE 2020*“. Jde zejména o přeshraniční informační služby v rámci eEnvironmentu. Jedná se o sdílení monitorovaných a zpracovávaných dat a informací o atmosféře, povrchových i podzemních vodách, odpadech, půdě, biodiverzitě, atd. pomocí Globálního monitorovacího systému životního prostředí a bezpečnosti (GMES – Global Monitoring for Environment and Security). Tyto informační služby umožňují efektivnější a přesnější sledování aktuálního stavu životního prostředí a udržitelného rozvoje v Evropě, dále pak jeho modelování a simulaci jeho dalšího vývoje.

Za hlavní přínos 9. letní školy považujeme skutečnost, že na letošním ročníku setkali a jsou v příspěvcích zastoupeni nejen doktorandi a učitelé z Masarykovy university (Institut biostatistiky a analýz, Přírodovědecká fakulta a Fakulta informatiky), ale i z Mendlovy university (Provozně ekonomická fakulta), Vysokého učení technického (Podnikatelské fakulta) v Brně a Vysoké školy báňské – Technické university (Ekonomická fakulta) v Ostravě. Jejich příspěvky ve sborníku přispívají k tomu, že se letní škola stala širokým interdisciplinárním odborným fórem v rámci České republiky. Dále je důležité, že několik příspěvků, jejichž spoluautoři jsou ze zahraničí, je publikováno v anglickém jazyce, který podtrhuje mezinárodní význam sborníku.

Těžiště projednávaných otázek na letní škole bylo především v detailní diskusi věnované řešení projektu TaToo, ale zahrnulo i další problematiku, která se týkala oblasti „eEnvironment“, „eGovernment“, „eParticipation“ a prezentaci nově řešeného projektu podporovaného Grantovou agenturou České republiky s názvem „Construction of Methods for Multifactor Assessment of Company Complex Performance in Selected Sectors“, vedený pod registračním číslem P403/11/2085.

V Brně dne 31. listopadu 2012

Jiří Hřebíček
Jan Ministr
Tomáš Pitner

Editori

Obsah

Current Trends of Corporate Performance Evaluation and Reporting for Building and Construction Sector

Michal Hodinka, Ondřej Popelka, Jana Soukopová, Michael Štencl, Oldřich Trenz

Maple – pro marketing a inovace – nástroj měření výkonnosti podniku

Zuzana Chvátalová, Pavlína Charvátová

Evoluční model pesimizme prostředí v Maple

Jiří Kalina

Systém pro optimalizaci medicínského kurikula

Martin Komenda

Verification of TaToo tools from the perspective of Validation Scenarios Solved by Masaryk University Team

Miroslav Kubásek, Jiří Hřebíček

Tagging Tool based on a Semantic Discovery Framework – Semantic Framework Implementation

Miroslav Kubásek, Jiří Hřebíček, Sinan Yurtsever, Pascal Dihé, Sasa Nesic, Giuseppe Avellino, Luca Petronzio

Bezpečnost sociálních sítí na internetu

Jan Ministr

Scaling CEP to Infinity

Filip Nguyen, Tomáš Pitner

Monitorování a evaluace výukových procesů

Lucie Pekárková, Patrícia Eibenevová

Dohledové systémy

Tomáš Pitner

Ecosystem Condition Modeling Using Machine Learning Tools

Vadim Rukavitsyn

Indoor Navigation for Mobile Devices

Jonáš Ševčík

Princípy doručovania BI na mobilné zariadenia

Lucia Tokárová

Logování pro novou generaci monitoringu

Daniel Tovarňák, Tomáš Pitner

Computation of diffusion coefficients for lipids in polymers

Jaroslav Urbánek, Tatsiana P. Rusina, Foppe Smedes

Current Trends of Corporate Performance Evaluation and Reporting for Building and Construction Sector

Michal Hodinka¹, Ondřej Popelka¹, Jana Soukopová², Michael Štencel¹, Oldřich Trenz¹

¹Department of Informatics, Faculty of Business and Economics, Mendel University
Zemědělská 1, 61300 Brno, Czech Republic
{michal.hodinka, ondrej.popelka, michael.stencel, oldrich.trenz}@mendelu.cz

²Department of Public Economics, Faculty of Economics and Administration, Masaryk University
Lipová 41a, 602 00 Brno, Czech Republic
soukopova@econ.muni.cz

Abstract

Current trends of sustainability indicators evaluation (i.e. measurement of economic/financial, environmental, social and governance (ESG) performance) and corporate sustainable reporting are discussed in the paper. The focus is on the building and construction sector. The relationship between sustainability indicators and reporting is an important issue; and the development of advanced methods to identify key performance indicators for ESG performance is discussed here along with the possibility of the utilization of information and communication technology and XBRL taxonomy.

Abstrakt

V příspěvku jsou popsány současné trendy v hodnocení indikátorů udržitelnosti (tj. měření ekonomické / finanční, environmentální, sociální a správní (ESG) výkonnosti) a podnikového reportingu udržitelnosti. Příspěvek je zaměřen na stavebnictví a výstavbu. Vztah mezi indikátory udržitelnosti a reportingem je důležitou otázkou, a vývoj pokročilých metod k identifikaci klíčových výkonnostních indikátorů pro ESG výkonnosti je diskutován spolu s možností využití informačních a komunikačních technologií a XBRL taxonomie.

Key words

Performance evaluation, Corporate performance, Key performance indicators, Corporate sustainability reporting, GRI, UN Global Compact, UNEP FI, ISO 26000, XBRL, Building and construction sector

Klíčová slova

Hodnocení výkonnosti, firemní výkonnost, klíčové ukazatele výkonnosti, podnikový reporting udržitelnosti, GRI, UN Global Compact, UNEP FI, ISO 26000, XBRL, stavebnictví a výstavba

1. Introduction

Successful corporate sustainability, i.e., the capacity of an organization to continue operating over a long period of time, depends on the sustainability of its stakeholder relationships. The available statistics show that through all objective benefits the sustainability and ESG indicators evaluation and corporate sustainability reporting can bring an appropriate feedback to businesses. This research in the area of corporate performance evaluation and corporate sustainability reporting [13], [14], [15], [16], [24] and [23] reflects the overall global world trends [1], [7], [31], [33], [35].

We have analysed corporate performance and ESG factors in chosen companies of the construction and real estate sector which have implemented and certified international management standards [25], i.e. quality (ISO 9000), environmental (ISO 14000 and EMAS) and occupational health and safety (ISO 18000) management systems, and some of them are going to implement the corporate social responsibility (ISO 26000) management system.

Therefore, ESG data and information are being monitored, codified, registered and transformed to Key Performance Indicators (KPIs) [1], [8], [12], [14], [15], [16] and [31]. This fact indirectly indicates that, in the case of such needs, the organization is able to add this ESG data and incorporate it into the corporate sustainability report, [5], [8], [15], [16].

In the paper we summarize chosen results of project No P403/11/1103 of the analysis of ESG aspects of corporate performance evaluation and reporting issued by the Global Reporting Initiative (GRI) which provides Sector Supplement for all reporting organizations in the construction and real estate sector [4].

We have focused on the critical partial processes in our research areas: *integration of economic, environmental, social and governance performance*.

Our analyses of possibilities of corporate performance measurements in chosen organizations of the construction and real estate sector by means KPIs were based on analyses of previous findings [13], [14], [15], [16], [24] and their results will also be discussed in the paper.

2. New approach of GRI reporting

In this chapter we introduce some results of our analysis of the state-of-the-art economic, environmental, social and governance aspects of corporate performance of the construction and real estate sector. There we focused on the new approach of GRI reporting developed with other organizations on common approaches to corporate performance and reporting [13], [15] and considered the important European Union (EU) legislation (i.e. *Construction Product Regulation (CPR)* - Regulation (EU) No 305/2011 laying down harmonized conditions for the marketing of construction products) and the *United Nations Environment Programme Sustainable Buildings and Climate Initiative (UNEP SBCI)* [38].

The *Global Reporting Initiative* is a very important network-based organization that produces a comprehensive sustainability reporting framework that is widely used around the world. The GRI has pioneered the development of the world's most widely used sustainability reporting framework in 2000 and is committed to its continuous improvement and application worldwide. The GRI drives sustainability reporting by all organizations. It produces the world's most comprehensive *Sustainability Reporting Framework* (GRI Framework) [34] which is the family of reporting guidance materials provided by GRI.

Sustainability reports based on the GRI Framework can be used to demonstrate an organizational commitment to sustainable development, to compare organizational performance over time, and to measure organizational performance with respect to laws, norms, standards and voluntary initiatives.

GRI's Framework consists of the *Sustainability Reporting Guidelines*, Sector Guidance's, National Annexes, and the Boundary and Technical Protocols [6].

The GRI promotes a standardized approach to reporting to stimulate demand for information on sustainability – benefitting both reporting organizations and report users.

In March 2011, the GRI released the G3.1 Guidelines [7], an update and completion of the G3 Guidelines from 2006 [6], which consists of two parts. Part 1 features guidance on how to report. Part 2 features guidance on what should be reported. This is defined in the form of *Disclosures on Management Approach (DMA)* and *Performance Indicators (PI)*, which are organized into categories: *Economic, Environmental and Social*. The Social category is broken down further to Labor, Human Rights, Society and Product Responsibility subcategories. Each category includes a DMA and a corresponding set of *Core and Additional Performance Indicators*.

Core Performance Indicators (CPI) have been developed through GRI's multi-stakeholder processes, which are intended to identify generally applicable PIs and are assumed to be very important for most organizations. An organization should report on CPIs unless they are deemed not material on the basis of the GRI Reporting Principles. Further we will take into account that CPIs can be in compliance with KPIs.

Additional Performance Indicators (API) represent emerging practice or address topics that may be material for some organizations, but are not material for others. Further we will not take into account APIs and try to identify only CPIs or KPIs.

The DMA should provide a brief overview of the organization's management approach to the *Aspects* defined under each *Indicator Category* in order to set the context for performance information. The organization can structure its DMA to cover the full range of *Aspects* under a given *Category* or it can group its responses on the *Aspects* differently. However, the DMA should address all of the *Aspects* associated with each category regardless of the format or grouping.

GRI PIs are first organized by a general sustainability *Category* (economic, environmental, social: labor; human rights; society; product responsibility), and then they are further arranged under *Aspect* headings which more specifically reflect the issue each indicator is designed to measure.

Although the G3.1 Guidelines [7] has served as an essential and very useful tools in improving the standardization of organization's reporting in many sectors, organizations continue to have differing degrees of compliance with the G3.1 Guidelines and sometimes also differing views on the best tools to apply these standards to their reporting. The integration of financial performance within environmental, social and governance performance reflects a growing desire by stakeholders for more information on a broader range of issues. To be comparable across all organizations, and thus useful for mainstream investment analyses, it is important that financial, environmental, social and governance (ESG) data are transformed into consistent units and presented in a balanced and coherent manner in *ESG indicators* [8].

G4 Guidelines is coming GRI's fourth generation of Sustainability Reporting Guidelines and is now in development.

The main focus of G4 Guidelines is:

- a general revision to improve DMA and PIs technical definitions;
- an extra effort to harmonize with other relevant international reporting guidance, see for example [13], [15];
- a considerably improvement of guidance around the definition of what is material (from different perspectives);
- a re-design of the G4 Guidelines format (by separating "standard like" requests from guidance, making it web based, offering templates, linking it to technology solutions using XBRL taxonomy).

The launch of the fourth generation of G4 Guidelines is planned for 2013. They will be developed using the international multi-stakeholder consultation process. Open Public Comment Periods, diverse expert Working Groups and GRI's approval procedures will ensure that G4's guidance will be in consensus based and reflects the broadest possible stakeholder input.

2.1 Guidelines for Construction and Real Estate Sector

The *Construction and Real Estate Supplement* (CRESS) provides organizations in the sector with a tailored version of GRI's Reporting Guidelines. It includes the original Guidelines, which set out the Reporting Principles, DMA and PIs for economic, environmental and social issues.

The CRESS is intended for companies that:

- invest in, develop, construct, or manage buildings; and
- invest in, develop or construct infrastructure.

The lifecycle diagram below describes the activity areas covered within the CRESS:

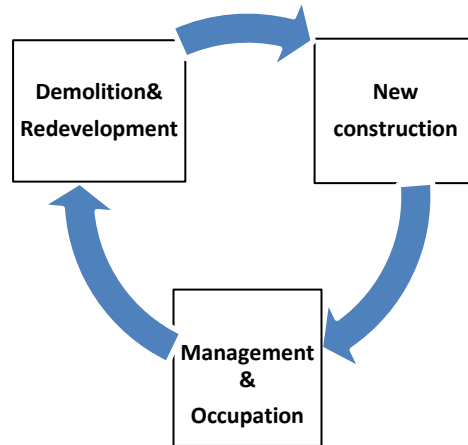


Figure 1: Lifecycle diagram of activity areas covered within CRESS. Source [4]

The construction and real estate sector has a significant impact on the economy, society, and environment, in ways that are both positive and negative.

The UNEP SBCI [38] suggests that buildings are responsible for more than 40 % of global energy use and one third of global greenhouse gas emissions. It also estimates that buildings are responsible for up to 80 % of greenhouse gas emissions in our cities and towns. Reducing global greenhouse gas emissions in the built environment is also widely recognized as the least expensive abatement opportunity. GRI recognizes that the construction and real estate sector has a significant role to play in the response to climate change.

Activities associated with constructing, operating, occupying and demolishing buildings and infrastructure also deplete natural resources and contribute many kinds of pollutant to land, air and water. Resources which are vital to the survival of all species, such as water and natural materials, are consumed on a significant scale by activities associated with the built environment. The UNEP SBCI estimates that the built environment is globally responsible for 30 % of natural material use and 20 % of water use.

The creation and maintenance of the built environment also significantly affects natural ecosystems and transforms or eradicates long standing habitats. The construction and real estate sector also produces large quantities of waste and UNEP SBCI estimates that the built environment contributes to 30% of total solid waste generation.

In socio-economic terms, the built environment has significant direct and indirect impacts on social wellbeing and the livelihoods and prosperity of communities and individuals. The sector, through its various activities as a major employer with a diverse and complex supply chain, can positively impact local economies by providing jobs, training and industry. The sector provides homes, education and recreational facilities for communities, yet it can also be responsible for displacing many people.

The sector's products are also enduring, in some instances lasting hundreds of years and forever changing the landscape in which they sit. These reasons, combined with the growing appetite for sustainability information from stakeholders and an increasing number of companies managing and reporting on their performance, have given rise to the need for reporting guidance and this Sector Supplement.

2.2 GRI and ISO reporting

The *International Organization for Standardization (ISO)* [20], the world's largest developer of voluntary International Standards, and the GRI, signed a *Memorandum of Understanding (MoU)* on 5 September 2011 to increase their cooperation. The MoU is intended to leverage the activities of the two organizations related to reporting and benchmarking by businesses and on sustainable development by sharing information on ISO standards and GRI programs, teaming up with other

partners, participating in the development of new or revised documents, joint promotion and communication. ISO and GRI are also meant to support and promote each other's involvement in initiatives related to sustainable development, such as the Rio+20 conference in Brazil in 2012, and other programmes by organizations such as the *United Nations Global Compact* [39], the *Organization for Economic Co-operation and Development (OECD)* [26], and the *United Nations Environment Programme Finance Initiative (UNEP FI)*, [37]. The ISO 26000:2010 *Guidance Standard on Social Responsibility* emphasizes the value of public reporting on social responsibility performance for internal and external stakeholders, such as employees, local communities, investors and regulators. ISO 26000 provides guidance on the underlying principles of social responsibility, the core subjects and issues pertaining to social responsibility and on ways to integrate socially responsible behaviour into existing organizational strategies, systems, practices and processes. ISO 26000 also emphasizes the importance of results and improvements in social performance.

ISO 26000 also briefly explains that social responsibility reports and other communications should be understandable, accurate, balanced/transparent, and timely, as well as comparable. The GRI Framework goes further in providing more specific guidance on the principles of clarity, accuracy, balance, timeliness, and comparability and also adds the principle of reliability. These principles all go towards helping to ensure the quality of reported information.

This represents an important new level of international attention with respect to the issue of reporting, and is aligned with GRI's vision that disclosure on economic, environmental, social and governance performance becomes as common place and comparable as financial reporting [9].

2.3 GRI and Integrated Reporting

The *International Integrated Reporting Council (IIRC)* [18] was established to support the evolution of integrated reporting. The IIRC brings together the world's leaders from the corporate, investment, accounting, securities, regulatory, academic and standard-setting sectors, as well as civil society. The IIRC aims to develop a new approach to reporting – one that is fit for purpose in the 21st Century – building on the foundations of financial, narrative, governance and sustainability reporting, but in a way that reflects the reality that all these elements are closely related and interdependent, and flow from the organization's overall strategy and business model. In September 2011 the IIRC published its discussion paper *Towards Integrated Reporting - Communicating Value in the 21st Century*, [36], which offers initial proposals for the development of an International Integrated Reporting Framework and outlines the next steps towards its creation and adoption.

GRI is one of the co-conveners of the IIRC and is actively participating in its working groups and task forces. GRI works towards making disclosure of sustainability impacts a mainstream business activity. There are different paths to mainstreaming, and many uses for corporate sustainability reporting: as a standalone discipline; as part of a company's research and development; as a platform for providing data to specific stakeholder groups, like investors; and now, as an intrinsic element of integrated reporting.

Integrated reporting is a form of corporate reporting that brings together material information about an organization's strategy, governance, performance and prospects in a way that reflects the commercial, political, social and environmental context within which it operates. It provides a clear and concise representation of how an organization creates value, now and in the future.

GRI supports the development of integrated reporting as it has the potential to make a large contribution to the mainstreaming disclosure of sustainability impacts.

2.4 GRI and the Carbon Disclosure reporting

GRI and the *Carbon Disclosure Project (CDP)* [2] announced in July 2011 the release of *Linking GRI and CDP: How are the GRI Guidelines and the CDP questions aligned?* The first edition of this document [24] was published in 2010 and has now been updated to incorporate changes in guidance. *Linking GRI and CDP* features a table that compares specific environmental indicators from GRI's Guidelines with questions from CDP's Investor and *CDP Supply Chain 2011 programs* [3].

2.5 GRI and XBRL Reporting

The eXtensible Business Reporting Language (XBRL) [42] is a markup language for the electronic communication of business and financial data that provides major benefits in the preparation, analysis and communication of business information.

XBRL is the emerging standard used around the world to define and exchange financial performance data. With substantial work already initiated internationally to create taxonomies for financial information, the GRI has, in collaboration with partners, developed XBRL taxonomy for non-financial performance data that can complement other taxonomies.

The GRI Taxonomy Project announced in June 2011 by GRI and Deloitte will result in a new format for exchanging sustainability data: *one that will help investors, auditors and analysts to publish, use and analyse information in sustainability reports more quickly and easily.*

This project will develop the XBRL taxonomy for GRI's G3 and G3.1 Guidelines, and is now underway.

The GRI Framework with the new XBRL taxonomy designed for the ESG performance together with renewed EU strategy 2011-14 for *Corporate Social Responsibility (CSR)* [30], *UN Compact Global*, *UNEP FI*, *ISO 26000* and *OECD Guidelines for Multinational Enterprises* [27] appear as essential for corporate reporting at present.

2.6 Summary of GRI reporting

Data on the corporate performance including carbon emissions, water use and human rights infringements can now be easily revealed thanks to a new format for tagging data in sustainability reports, being launched on 8 March 2012 by the GRI. This new format will help people find information hidden in corporate sustainability reports much more quickly and easily.

According to recent research, 95 % of the world's 250 biggest companies now report their sustainability performance. GRI produces a comprehensive sustainability reporting framework that is widely used around the world. The Framework, which includes the Sustainability Reporting Guidelines, features indicators that organizations can use to measure and report their sustainability performance.

In the past decade, corporate reporting has evolved to include sustainability information, on the economic, social and environmental performance of an organization. Around the world, more companies are releasing sustainability performance information, both through annual sustainability reports or an equivalent document, and – because of the increasing demand for it – also through other means, such as websites, newsletters and other corporate reports. Increasingly, companies are integrating sustainability disclosures into their regular reporting cycle. Today, some 4,500 organizations report their sustainability performance.

GRI has launched a new XBRL taxonomy for tagging sustainability data in reports, making it easier for report users – including regulators, investors and analysts – to find and analyse data. The GRI Taxonomy – which is available for free – was developed in collaboration with Deloitte Netherlands. A team of experts from different stakeholder communities reviewed the draft taxonomy before the Public Comment Period.

GRI taxonomy will enable companies and other organizations to use XBRL to improve their sustainability reporting and make the data in their reports more accessible.

Nelmara Arbex, Deputy Chief Executive of the GRI, said: *“Today's new taxonomy is a major step forward in making sustainability data available to society. Many companies already use XBRL to tag their financial performance data; the GRI Taxonomy means that companies can tag their sustainability data, making it easily accessible for people who want to find information in the report.”*

Tagging ESG data in reports requires a piece of software. Some regulators – including stock exchanges and governments – use various ICT tools to search for data and compare the performance of different companies.

We are going to develop ICT tools in the project No P403/11/1103 for corporate sustainability reporting for Construction and Real Estate Sector.

3. Corporate performance evaluation and reporting

The corporate performance plays a key role in the corporate strategic policy and sustainability of success of an organization. The creation of reliable methods of ESG performance measurement where concurrent acting of multiple factors is in play can be considered a prerequisite for success not only in decision-making, but also with regard to corporate governance, comparison possibilities, development of a healthy competition environment etc.

The GRI Framework states that corporate performance indicators may be both quantitative and qualitative and that they should cover the reporting entity's direct and indirect impacts across economic, environmental and social dimensions.

Economic indicators include proxies for the organization's impact on resources at the shareholder level and on other economic systems at the local, national and global level. This heading also encompasses issues dealing with remuneration paid to employees and money received from customers, to name but a few.

Environmental indicators deal with the measurement of an organization's impact on the environment via its products and services and its activities.

Social indicators deal with labor practices, human rights and broader social issues affecting a broad range of stakeholders [43]. An important element of the social performance is occupational health and safety. The trend underscoring the social aspects of sustainable development is the concept of CSR [44]. Other key issues related to the CSR are: human rights, employees' rights, involvement of municipalities and relationships with suppliers, information policy including issues such as releasing information, transparency, educating the consumers and anti-corruption measures.

Governance indicators enlarge *Sustainability indicators* and deal with corporate governance. This is a term that refers broadly to the rules, processes, or laws by which businesses are operated, regulated, and controlled. The term can refer to internal governance indicators/factors defined by the officers, stockholders or constitution of a corporation, as well as to external forces such as consumer groups, clients, and government regulations. The corporate governance issues in the Czech Republic are obtained from the *Corporate Governance Code* of companies, which is based on the OECD principles 2004 [27].

One of the possible approaches is to also take into account successful solutions to economic, environmental and social issues and governance in relation to measurement of corporate performance, as well as its continued success (Sustainability of Success). Disregarding such aspects of performance in the unified reporting (e.g. prepared G4 Guidelines for Corporate Sustainability Reporting) by company managers may result in creating further and even deeper problems. For the purpose of collecting corporate performance data it is necessary to determine the KPIs of the given organization.

3.1 Integration of economic performance

We will here consider economic performance based on the G3.1 Guideline (enlarged with its Construction and Real Estate Supplement). Economic performance indicators are often used for selection strategies (maximizing profits, maximizing total costs, company survival, etc.) based on direct economic impacts of customers, suppliers, employees, providers of capital, public etc.

Financial reporting standards, such as IFRS and US Generally Accepted Accounting Principles (U.S. GAAP) and ESG reporting frameworks, principally the GRI Guidelines [7, 8], will act as structural supports for potential integrated reporting frameworks of integrated economic performance [17].

Research of the direction of the economic performance indicators of project No P403/11/2085 has focused on the analysis of the reporting framework of the GRI [7] and IFAC Sustainability Framework 2.0 [33]. Furthermore, the research dealt with economic indicators which have been published in the

Yearbook of Czech Statistical Office [45] and selected economic indicators of financial statements according to Czech accounting standards (from 2011) and a comprehensive analysis of the voluntary reporting of 10 large Czech companies of the Construction and Manufacturing sector has also been done [46].

We proposed the Key Performance Indicators (KPIs) for the measurement of economic performance in relation to the sustainability and ESG indicators. The economic performance indicators provide quantitative forms of feedback which reflect the results in the framework of corporate strategy. The approach is not different when we control environmental, social and governance issues. The non-financial KPIs that an organization develops, manages and ultimately reports – whether internally or externally – will depend on its strategic priorities, and will reflect the unique nature of the organization. What is most important is to recognize what is measured, what is controlled, and it is important that the measures create value for the company and its stakeholders.

The proposed KPIs can help organizations to plan and manage their economic priorities, in particular, when the economic indicators are focused on the core business strategy, by means of operational plans, which include performance targets.

Table 1. Economic KPIs. Source [46]

Indicator	KPIs	Measurement
EC1 Profit	EBIT	Earnings before Interest and Taxes
	EBITDA	Earnings before Interest, Taxes, Depreciation and Amortization.
	EAT	Earnings after Taxes / Net profit
	EPS	Earnings Per Share, P/E = Price Earnings Ratio.
EC2 Cash Flow	FCF Free Cash Flow	EBIT * (1-Tax rate) + Depreciation and Amortization - Changes in Working Capital - Capital expenditure.
	OCF Operating Cash Flow	All the cash flows arising from the main activity of the company, which is the subject of its business (the movement of stocks, receivables, obligations).
EC3 Revenues	TR Total revenues	Total revenue is the total receipts of a company from the sale of any given quantity of a product, i.e. Revenues from own goods and services + Revenues from sale of merchandise (goods for resale) + Revenues of fixed assets + Revenues from sale of materials + Revenues of securities.
EC4 Turnover size	Turnover size	Revenues from own goods and services + Revenues from sale of merchandise (goods for resale) + Revenues of securities
EC5 Profit margin	Profit margin	The difference between turnover (revenues) from sales of goods and expenses on merchandise sold (i.e. on goods sold in the same condition as received).
EC6 Indicators of economic performance	Return on Equity	ROE = EAT / Equity
	Return on Investment	ROI = EBIT / Total capital
	Return on Assets	ROA = EBIT / Assets
	Return on Sales	ROS = EAT / Revenues
	Return On Capital Employed	ROCE = EBIT / Equity + Long-term liabilities
EC7 EVA	Economic Value Added	EVA = (ROE – Cost of Equity) * Equity

The proposed KPIs for measurement of the corporate performance in relation to the ESG indicators were established on the basis of the results of empirical research by the team of FBM BUT, [46], see Tab. 1.

These indicators EC1 – EC7 differ from indicators proposed in CRESS [4, 30], where they are defined only in general following GRI 3.0 Guidelines:

- *Economic Performance indicators*: EC1 (Commentary added to clarify sources of financial information. Commentary added to report on specific breakdown for payments to governments. Commentary added to refer to methodology for calculating community investments and clarifying infrastructure investments) and EC2 (Commentary added to report financial implications and other risks and opportunities for the organization's activities due to other sustainability issues. Commentary added to provide new definitions on Qualitative Financial implications and Obsolescence).
- *Market Presence indicator*: EC7 (Commentary added to include procedures for local hiring for all direct employees, contractors and sub-contractors hired from the local community. Commentary added to provide definitions on contractors and sub-contractors).
- *Indirect Economic Impact indicators*: EC8 (Commentary added to explain other significant infrastructure investments made by the reporting organization.) and EC9 (Commentary added to add examples of indirect economic impacts).

All economic performance indicators EC1 – EC7 of Table 1 give measureable values. A company in the Czech Republic can compare some of them with the country's benchmark value, e.g., it is able to calculate the EVA indicator and compare with the benchmark value online on the web of the Ministry of Industry and Trade of the Czech Republic¹.

We took also into account GRI's Reporting Guidelines CRESS and consider following economic KPIs:

- EC1 - Direct economic value generated and distributed, including revenues, operating costs, employee compensation, donations and other community investments, retained earnings, and payments to capital providers and governments.
- EC2 - Financial implications and other risks and opportunities for the organization's activities due to climate change and other sustainability issues.
- EC7 - Procedures for local hiring and proportion of senior management and all direct employees, contractors and subcontractors hired from the local community at locations of significant operation.
- EC8 - Development and impact of infrastructure investments and services provided primarily for public benefit through commercial, in kind, or pro bono engagement.
- EC9 - Understanding and describing significant indirect economic impacts, including the extent of impacts.

Financial reporting standards, such as the *International Financial Reporting Standards (IFRS)* [17] and the *US Generally Accepted Accounting Principles (U.S. GAAP)* [40] and ESG reporting frameworks, principally the GRI Guidelines and our proposed set of economic indicators, will act as structural supports for potential integrated reporting frameworks of integrated economic performance. The IFRS Foundation, the body that oversees the *International Accounting Standards Boards (IASB)*, has today completed the first part of their project to address requests by regulators and preparers for extensions to the full *IFRS XBRL Taxonomy*, which we are going to use in our developed ICT tools for reporting.

The IFRS XBRL Taxonomy is used to help those filing IFRS financial statements electronically to “tag” the information with identification tags (called “concepts” in an XBRL taxonomy). Currently, the IFRS XBRL taxonomy includes all core concepts included in IFRSs as issued by the IASB.

We have also used our developed XBRL tools to facilitate the calculations and the visualizations of these integrated economic performance indicators [47].

¹ <http://www.mpo.cz/cz/infa.html>

3.2 Integration of environmental performance

We determined general environmental KPIs with the use of results of our previous research in this field [11], [14] and with the use of the G3.1 Guideline and EMAS indicators for all sectors, which were accepted by the Ministry of Environment of the Czech Republic as its official methodology for a voluntary environmental reporting [12]. The proposed environmental KPIs shall apply to all organizations in all NACE economic activity sectors.

We have identified direct and indirect environmental aspects of construction sectors, where we issued Reference Documents on Best Environmental Management Practice in the Building and Construction Sector [48], see Fig. 2.

However we have selected a certain set of environmental KPIs following key areas of the environment from GRI's Reporting Guidelines CRESS [30] and EMAS and used GRI notations:

- 1) *Efficiency of material consumption*, where we have chosen EN1 and EN2 indicators from CRESS;
- 2) *Energy efficiency*, where we have selected EN3, EN4, EN5, EN6, EN7 indicators and an additional CRE1 indicator (Building Energy Intensity) from CRESS;
- 3) *Water management*, where we have selected EN8, EN9, EN10 indicators and additional CRE2 indicator (Building Water Intensity) from CRESS;
- 4) *Waste management*, where we have selected EN22 indicator from CERSS and additional EN22a indicator (Total annual generation of hazardous waste) from [1];
- 5) *Biodiversity*, where we have selected EN12 and EN13 indicators from CRESS;
- 6) *Air pollution*, where we have selected EN16, EN17, EN18, EN20 indicators and additional CRE3 indicators (Greenhouse gas emissions intensity from buildings) and CRE4 (Greenhouse gas emissions intensity from new construction and redevelopment activity) from CRESS;
- 7) *Other relevant indicators of the influence of the organization's activity on the environment*, where we have selected EN26, EN29 indicators and additional CRE5 indicator (Land and other assets remediated and in need of remediation for the existing or intended land use according to applicable legal designations) from CRESS.

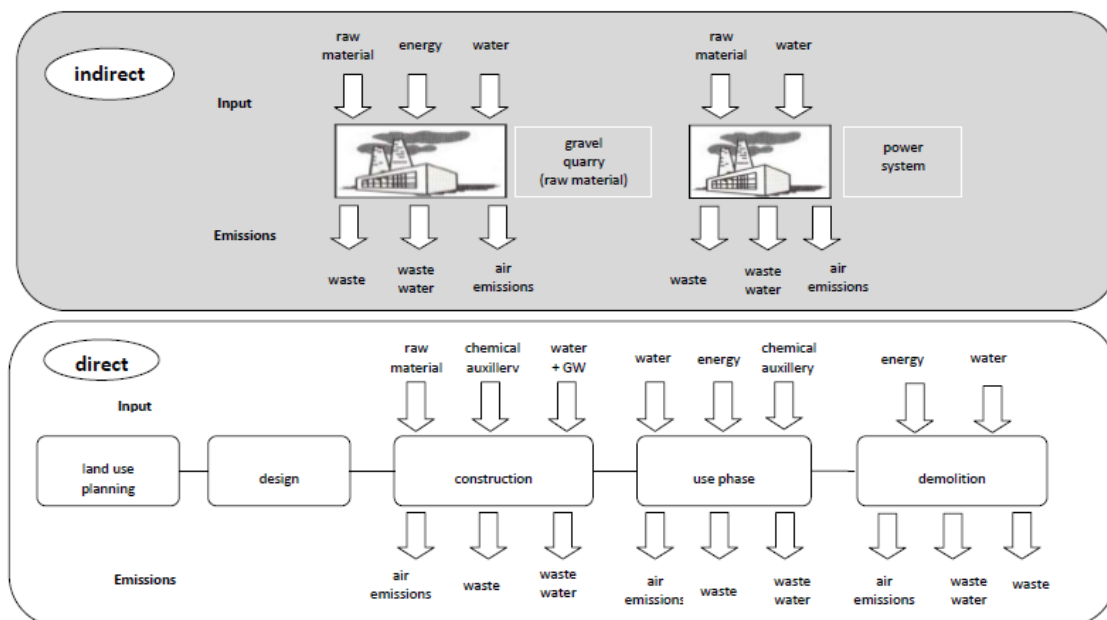


Figure 2: Direct and indirect environmental aspects of the construction sector. Source: [48]

The above set of selected environmental KPIs differs from our past set of KPIs introduced in [11], [12], [13] and describes more appropriate KPIs in sustainability and ESG indicators for the building and construction sector.

Some constructions of KPIs represent absolute performance (e.g., total GHG emissions, total water use), which is not normalized by factors such as floor area or building users. However, where it is practical to do so and will be helpful in interpretation, the reporting organizations should consider using ‘like-for-like’ analysis for absolute KPIs to enable comparability over a defined period of time of our research project.

3.3 Integration of social performance

The social dimension of corporate sustainability concerns the impacts the given organization has on the social systems within which it operates. We are going to determine the KPIs for social performance based on the GRI Framework and its social performance indicators, in order to identify some key performance aspects surrounding *labour practices*, *human rights*, *society*, and *product responsibility* [7], [14], as was done in GRI’s Reporting Guidelines CRESS.

We have to consider that *labour practices indicators* also draw upon two instruments which directly address the social responsibilities of business enterprises: the *ILO Tripartite Declaration Concerning Multinational Enterprises and Social Policy* [19], and the *OECD Guidelines for Multinational Enterprises* [27] and we must take into account: employment; labour/management relations; health and safety; training and education; diversity and opportunity.

However we are going to select the optimal set of social KPIs for NACE economic activities: “F – Construction” in the following key areas:

- 1) *Labor Practices and Decent Work indicators* are broadly based on the concept of decent work. The set begins with disclosures on the scope and diversity of the reporting organization’s workforce, emphasizing aspects of gender and age distribution. We here take into account:
 - *Employment* - LA1 and LA3 indicators from CRESS;
 - *Occupational Health and Safety* - LA7, LA8 and CRE6 (Percentage of the organization operating in verified compliance with an internationally recognized health and safety management system) indicators from CRESS;
 - *Training and Education* - LA10 indicator from CRESS;
 - *Diversity and Equal Opportunity* - LA13 indicator from CRESS;
 - *Equal Remuneration for Women and Men* - LA14 indicator from CRESS;
- 2) *Human Rights indicators* require companies to report on the extent to which human rights are considered in investment and supplier/contractor selection practices. We here take into account:
 - *Non-discrimination* - HR4 indicator from CRESS;;
 - *Child labour* - HR6 indicator from CRESS;
- 3) *Society indicators* focus the attention on the impacts organizations have on the communities in which they operate, and disclosing how the risks that may arise from interactions with other social institutions are managed and mediated. In particular, information on the risks associated with bribery and corruption is sought, as well as information on the undue influence in public policy-making, and monopoly practices. We here take into account:
 - *Local community* – SO1, SO9 and CRE7 (Number of persons voluntarily and involuntarily displaced and/or resettled by development, broken down by project) indicators from CRESS;
 - *Public policy* – SO5 and SO6 indicators from CRESS.
- 4) *Product responsibility indicators* address the aspects of a reporting organization’s products and services that directly affect customers. We take into account namely:
 - *Customer Health and Safety* – PR1 and PR2 indicators from CRESS;;
 - *Products and Services Labelling* – PR3, PR4, PR5 and CRE8 (Type and number of sustainability certification, rating and labelling schemes for new construction, management, occupation and redevelopment) indicators from CRESS.

The integration process of the development of the complete set of social performance indicators is in progress and the final version of KPIs is planned to be complete, as a part of our research project, towards the end of this year.

Certain KPIs should also be reported by meaningful segmentation to facilitate interpretation, for example by portfolio, fund, geographic location, or asset type.

3.4 Integration of corporate governance performance

The corporate sustainability or ESG reporting usually contains governance structure of the organization, including committees under the highest governance body responsible for specific tasks, such as setting the strategy or organizational oversight (CEO, top management etc.).

Legal base for the corporate governance is created within the framework of following EU directives and rules (Code of the Criminal Responsibility, Code of the Business Activities on the Financial Markets, Commercial Law, Principles of Auditors, and Bank Law) and from others legislative (Directive 2004/25/ES about offers undertaking), about the transparency of the listed corporation (Directive 2004/109/ES), right of shareholders (Directive 2007/36/ES), about market exploitation (Directive 2003/6/ES) and about the audit (Directive 2006/43/ES).

The corporate governance regulation in the Czech Republic usually uses a dualistic model: the mechanism of written law enforcement (mainly the Act No 513/1991 Sb., Commercial Code), and the self-regulation mechanism, characterized by a self-imposed observance of the required rules. This mechanism is primarily implemented through the code of company governance and also through due diligence principles. The company is governed by a body of shareholders – the general meeting reported to by the board of directors as an executive managing body and by the supervisory board as a surveillance authority.

The establishment of corporate governance performance indicators was based on the empirical analysis of the Code of corporate governance of OECD [27] and the Czech Republic (2004); also on the „Green Paper“ of the EU Corporate Governance Framework [49] and International Federation of Accountants (IFAC) [50].

The results of the current monitoring of the indicators in the area of corporate governance are shown in the Table 2, where each indicator is presented with its description and source.

Table 2. List of Selected KPIs of the Corporate Governance. Source [51]

Title	Description	Source
Management	Frequency of the executive body sessions	Corporate Governance and Management Code
Ownership concentration	Concentration of owners – right to vote per models	The OECD principles of CG, Annex: Indicators of Corporate Landscape OECD 2007
	Percentage distribution of the ownership per various categories of the investors	
Members of the board	Number of members from the point of the professional competences	Green Book - a governance and management of the company in the financial institutions and remuneration policy
	Percentage representation from the point of the international representation	
	Percentage representation of the members from the point of both sexes	
Stakeholder effectiveness	Percentage representation of the independent members	Corporate Governance and Management Code
	Separation of the posts CEO/chairman	IFAC
	Independency of the board members and audit bodies	Corporate Governance and Management Code, IFAC
	Duration of the membership in the board	IFAC
	Remuneration of the board -stimuli	IFAC
	Remuneration of the board – quantity of bonuses	Recommendation of the Council 2009/385/EC dated April 30, 2009, amending the recommendations 2004/913/EC and 2005/162/EC
	Remuneration of the board – offer (purchase, sale of shares)	
	Remuneration of the board – quantity of shares versus salary	IFAC

	Remuneration of the board - long term and short term obstacles	
	Percentage of women in the board	
	Signs of the risk management and policy implementation- division of competencies for the risk management	
Stakeholder engagement	Frequency of the involvement of the stakeholders	IFAC
	Existence of the mechanisms of the involvement of the stakeholders	
	Methods of the responses to the feedback from the stakeholders	
Conduct, litigation risk corruption	Records on the breaching of the regulations and extra costs	IFAC
	Corruption in comparison to the percentage of revenues in the region	
	Corruption - number of the analyzed business units	
	Total sum spent on the correction, penalties, expenses and putting out of operation	
	Payments to the state and the total value of the financial and subsistence contributions to the political parties, politicians and allied institutions	
	Right of vote equality	

3.5 ICT tools for corporate reporting and project

We have also used our developed ICT tools based on GRI XBRL taxonomy to facilitate the calculations and the visualizations of key performance indicators.

XBRL allows us to prepare reports to place electronic tags on specific content (graphs, numbers, text, etc.) of indicators in their reports by using an existing “*XBRL taxonomy*”. It enables users who are interested in finding some environmental data, e.g. on greenhouse gas (GHG) emissions, to immediately find this data - select it, analyze it, store it and exchange it with other computers on an intranet or internet network and automatically present it in different ways. Users are also able to apply this in a variety of reports and to compare emissions information across different reports.

The power of XBRL comes from its structure, which is divided into an instance document and a group of taxonomies. The instant document includes business facts that are reported. The XBRL taxonomy is defined by metadata about reported facts, meanings, interconnections, etc. From a technical point of view, the XBRL taxonomy is defined as a standardized XML schema (XSD), one that includes a concept, the data of which will appear in the report. The group of schemas describes the syntax as interconnections of the individual messages or their parts. Any XBRL schema goes hand-in-hand with so-called link bases.

Link bases are collections of references, which enrich the syntax by means of certain semantics. The main distribution occurs in schemas and Link bases according to [28]:

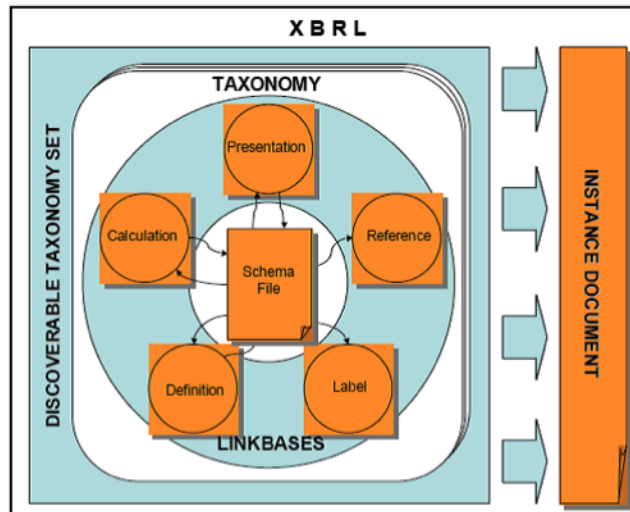


Figure 3: XBRL framework [28]

- 1) The core is the XML schema as the key unit of the XBRL taxonomy.
- 2) The LabelLinkBase, the naming of the elements included in the XSD schema.
- 3) DefinitionLinkBase, establishes the hierarchy and the organization of all the units appearing in the report.
- 4) The PresentationLinkBase, structuring and distribution of the concepts without changing the hierarchy defined by the DefinitionLinkBase.
- 5) The ReferenceLinkBase, enables the interconnection of the elements from the XML schema with other information, such as directives or comments, and – owing to that – simplifies the understanding of the whole structure.

Languages
Czech
English
Russian

Faculty of Business and Economics
Department of informatics

News
Basic information
Working Groups
Workshops and conferences
Technics

Search
Search this site:

Working groups

- SoNet
- Moebius
- GAČR 403
 - Time Schedule
 - Project Team
 - Publications
 - Related Sources
 - Contact
 - Internal Resources

Rychlé odkazy

- Projekt GAČR 403
- Zaměstnanci

Home

Project information
Tue, 08/02/2011 - 15:18 — prochazka

Construction of Methods for Multifactor Assessment of Company Complex Performance in Selected Sectors

Applicant: doc. Ing. Alena Kocmanová, Ph.D. (Faculty of Business and Management of the Brno University of Technology)
Co-Applicant: Prof. RNDr. Jiří Hřebíček, CSc. (Faculty of Business and Economics, Mendel University in Brno)
Project registration code: P403/11/1103
Project period: 3 years
Start date: 01/01/2011
Project information at Faculty of Business and Management (BUT)
Project information at Department of Informatics of FBE (MENDELU)

Abstract

Economic, environmental, social and governance factors form the core business strategy, as part of its daily operations, the challenge to succeed, defense against the threat, the incentive opportunities. Very important is their connection with the company performance assessment based on an evaluation (measurement) of a company complex performance, determined by these factors. Using advanced quantitative methods (optimization / stochastic / dynamic, etc.) in a measurement of complex company performance can be refined on qualitative aspects by capturing them with non-deterministic methods (uncertainty / data mining / imprecision / fuzzy logic, etc). The aim of the project is the construction methods for evaluation (measurement) of a complex company performance with sophisticated use of multi-factor metrics, creating methodology for selected sectors of economic activity that respects a wide range of modifiable use, which will be universally applicable in commercial sector, and their outcomes in the form of monograph and software, articles, and papers at conferences of international significance.

Motivation

The dynamic development of CSR and Reporting in the CR takes place in an era where investors want more than just a good product or an attractive sound of a brand. Moreover, the growing demand for and offer of the reporting corresponds with the rapid development of information society. We need to take into account market trends and issues affecting society, analysis of risks and opportunities. Relevant data to be used in global research in this field is very difficult to obtain (experiments are difficult to control or repeat), and the same is true of monitoring the impacts of comprehensive external effects. Significant are the levels of ESG performance and company overall performance. Strategically important activities may improve the corporate governance, sustainable activities and CSR, they create the intangible value of the company. Responsible sustainable development of corporate resources is the basis for the target company strategies. The main executive message clearly explains why the responsible and sustainable development of corporate resources is the basis for comprehensive

INFORMATIKA NA MENDELOVĚ UNIVERZITĚ

(Geografické specializace)

User login
Username:
Password:

[Request new password](#)

Visual Paradigm for UML

Project information
Project aims
Project outputs

Česky
English
Google™ Matně vyhledávání

Construction of Methods for Multifactor Assessment of Company Complex Performance in Selected Sectors

Project GACR P403/11/2085

Project information

Complete name of project	Construction of methods for multifactor assessment of company complex performance in selected sectors
Registration code of project	P403/11/2085
Project solution time	3 roky
Applicant	doc. Ing. Alena Kocmanová, Ph.D. Faculty of bussiness and management, Brno University of technology
Co-applicant	prof. RNDr. Jiří Hřebíček, CSc. Faculty of business and economics, Mendel University in Brno
Start date	1. 1. 2011
Finish date	31. 12. 2013

VYSOKÉ
UCENÍ
TECHNICKÉ
V BRNĚ

Provozně
ekonomická
fakulta

Mendelova
univerzita
v Brně

More information

- TEAM
- BASIS OF THE PROJECT
- PLANNED OUTPUTS
- PLANNED PUBLICATIONS
- CONTACTS

Figure 4: Web portal of the research project No P403/11/1103.
Source <http://gacr403.cz/en>

We have already prepared a web portal Fig. 4 with initial information about the planned integration of using an appropriate XBRL taxonomy, which will be differentiated and focused at particular target groups of users regarding the goals of our research project No P403/11/1103. Further, a group of ICT tools will be prepared to use at the same portal to validate and store differentiated XBRL reports.

3.6 Questionnaire for the investigation of corporate performance and reporting blueprint

Our research project No P403/11/1103 consists of partial research targets mentioned above. These targets are connected with the particular project stages. In the first stage, the state-of-the-art analysis we have developed, in collaboration with researchers of the FBM BUT, a questionnaire covering all four general topics (reporting is included across all the topics) of our research [13], [15], [16] was created. According to this, the questionnaire is divided into four independent modules focusing on partial aspects of business development, particularly in the environmental, social, economic, and corporate governance management subsystems.

The questionnaire was prepared for both the printed and the online version (with an identical text) and after all data collection will be completed, both data sets will be merged for further data processing. The online data collection will be done by means of the *Research Laboratory* (ReLa) questionnaire system, which has been developed as a research project of the Institute of Marketing and Trade of the FBE MENDELU in Brno [32].

Based on the research results of the questionnaire, it is possible to evaluate the current state and potential corporate performance of the investigated organizations of construction and real estate sector on environmental, social, economic and corporate governance levels.

Subsequently, we will continue in the verification of the correctness of our approaches and development of KPIs for corporate performance and corporate sustainability reporting, proposed for organizations of the investigated sectors of the Czech Republic and the European Union.

4. Conclusion

Analysis of the state-of-art on economic, environmental, social and corporate governance aspects of company performance and corporate sustainability reporting has been presented. The proposed set of abovementioned sustainability indicators for all companies in a given building and construction sector monitors to a much greater extent the development dynamics, as up to now [4]. CEO decision-making is based on a qualified assessment (measurement) of a situation determined at the same time by multiple indicators, primarily in their horizontal development [1], [16], [31], [52]. In pursuit of an outstanding information force, emphasis is currently being placed not only on the absolute data, but primarily on the changing data and the analyses of changes of these changes. That is, the dynamics of systems is the focus of attention. Vertical analyses that are applied adequately then add a further dimension to the conditions for decision making. These were carried out in the project No P403/11/2085. In this context other methods have to be discussed: logical and empirical methods, methods of qualitative and quantitative research such as statistical modeling, see [14], [35].

5. Acknowledgements

This paper is supported by the Czech Science Foundation. Name of the Project: Construction of Methods for Multifactor Assessment of Company Complex Performance in Selected Sectors. Reg. No P403/11/2085.

6. Bibliography

- [1] Bassen A, Kovacs A. M. (2008) Environmental, Social and Governance Key Performance Indicators from a Capital Market Perspective, *Zeitschrift für Wirtschafts und Unternehmensethik*, Vol. 9, pp. 182-192.
- [2] Carbon Disclosure Project website (2011). [Online]. Available: <https://www.cdproject.net/>
- [3] CDP Supply Chain 2011 programs (2011) on CDP. [Online]. Available: <https://www.cdproject.net/en-US/Programmes/Pages/CDP-Supply-Chain.aspx>.
- [4] Construction and real estate sector website (2012) on Globalreporting. [Online]. Available: <https://www.globalreporting.org/reporting/sector-guidance/construction-and-real-estate/Pages/default.aspx>.
- [5] Eccles R. G., Cheng B., Saltzman, D. (2011) The Landscape of Integrated Reporting. Reflections and Next Steps, Cambridge, Massachusetts, Harvard Business School, [Online]. Available: http://www.dvfa.de/files/finanzkommunikation/integrated_reporting/application/pdf/the_landscape_of_integrated_reporting.pdf.
- [6] G3 Guidelines website (2006) on Globalreporting. [Online]. Available: <https://www.globalreporting.org/reporting/latest-guidelines/g3-guidelines/Pages/default.aspx>.
- [7] G3.1 Guidelines (2011) on Globalreporting. [Online]. Available: <https://www.globalreporting.org/reporting/latest-guidelines/g3-1-guidelines/Pages/default.aspx>.
- [8] Garz H., Schnella F., Frank, R. (2010) KPIs for ESG. A Guideline for the Integration of ESG into Financial Analysis and Corporate Validation. Version 3.0, Frankfurt, DVFA/EFFAS, [Online]. Available: http://www.dvfa.de/files/die_dvfa/kommissionen/non_financials/application/pdf/KPIs_ESG_FINAL.pdf.
- [9] GRI and ISO 26000: How to use the GRI Guidelines in conjunction with ISO 26000. (2011), [Online]. Available: http://www.globalreporting.org/NR/rdonlyres/4A15F3C6-13D1-4AB4-9D5D-C93C38E56A3D/5468/ISOGRIReport_FINAL.pdf.

- [10] Hřebíček J., Soukopová J., Kutová E. (2010) Standardization of Key Performance Indicators for Environmental Management and Reporting in the Czech Republic. *International Journal of Energy and Environment*, Vol. 4, Issue 4, 2010, pp. 169-176.
- [11] Hřebíček J., Misařová P., Hyršlová J. (2007) Environmental Key Performance Indicators and Corporate Reporting, In *Proc. International conference EA-SDI 2007. Environmental Accounting and Sustainable Development Indicators*. Praha, Czech Republic: University Jana Evangelisty Purkyně, pp. 147-155.
- [12] Hřebíček J., Soukopová J., Kutová E. (2010) Methodological Guideline. Proposal of Indicators for Environmental Reporting and Annual Reports of EMAS. (in Czech)", Praha: Ministry of Environment of the Czech Republic.
- [13] Hřebíček J., Soukopová J., Štencl M., Trenz O. (2011) Corporate Performance Evaluation and Reporting", in *Proc. International Conference on Environment, Economics, Energy, Devices, Systems, Communications, Computers, Pure and Applied Mathematics*. Wiscontin, USA:WSEAS, pp. 338-343.
- [14] Hřebíček J., Soukopová J., Štencl M., Trenz O. (2011) Corporate Key Performance Indicators for Environmental Management and Reporting, *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis*, Vol. 59, pp. 99 - 108.
- [15] Hřebíček J., Soukopová J., Štencl M., Trenz O. (2011) Integration of Economic, Environmental, Social and Corporate Governance Performance and Reporting in Enterprises, *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis* Vol. 59, pp. 157—177.
- [16] Chvátalová Z., Kocmanová A., Dočekalová, M. (2011) Corporate Sustainability Reporting and Measuring Corporate Performance, In: *Proc. Environmental Software Systems. Frameworks of eEnvironment. 9th IFIP WG 5.11 International Symposium. ISESS 2011*. Heidelberg:Springer, pp. 398-406.
- [17] The IFRS website (2012). [online]. Available: <http://www.ifrs.org/>.
- [18] The International Integrated Reporting Council website. (2012) [Online]. Available: <http://www.theiirc.org/>
- [19] Tripartite Declaration Concerning Multinational Enterprises and Social Policy (2011) on ILO. [Online]. Available: http://www.ilo.org/wcmsp5/groups/public/---ed_emp/---emp_ent/---multi/documents/publication/wcms_094386.pdf
- [20] The ISO website (2012). [Online]. Available: <http://www.iso.org/iso/home.html>
- [21] Isenmann R., Gomez J. (2009) Advanced corporate sustainability reporting – XBRL taxonomy for sustainability reports based on the G3-guidelines of the Global Reporting Initiative“. In: *Proc. Towards eEnvironment. European Conference of the Czech Presidency of the Council of the EU 2009: Opportunities of SEIS and SISE: Integrating Environmental Knowledge in Europe*; March 25-27, 2009; Prague, Czech Republic. Prague, pp. 32-48.
- [22] Kocmanová A., Dočekalová M., Němeček P., Šimberová I. (2011) Sustainability: Environmental, Social and Corporate Governance Performance in Czech SMEs.” In: *Proc. the 15th World Multi-Conference on Systemics, Cybernetics and Informatics*. IFSR, Orlando, USA: WMSCI 2011, pp. 94-99.
- [23] Kocmanová A., Hornugová J., Klímková M., (2010) *Sustainability : The Integration Environmental, Social and Economic Performance of Company* (in Czech), Brno, Czech Republic: CERM.
- [24] Linking GRI and CDP: How are the Global Reporting Initiative Guidelines and the Carbon Disclosure Project questions aligned? (2011), [Online]. Available: http://www.globalreporting.org/NR/rdonlyres/4A15F3C6-13D1-4AB4-9D5D-C93C38E56A3D/6248/LinkingupGRIandCDP2011_ARcomments_VBeditMF.pdf.
- [25] Marimba A., Farinha J. T., Ferreira L. (2010) International standards integration for ecologic asset management“, In: *EE'10 Proceedings of the 5th IASME/WSEAS international conference on Energy & environment*, Wisconsin, WSEAS, pp. 44-51.
- [26] The OECD website. (2012) [Online]. Available: <http://www.oecd.org/>
- [27] OECD Guidelines for Multinational Enterprises (2011) on OECD. [Online]. Available: <http://www.oecd.org/dataoecd/43/29/48004323.pdf>.

- [28] Perrini F., Tencati A. (2006) Sustainability and Stakeholder Management: the Need for New Corporate Performance Evaluation and Reporting Systems. *Business Strategy and the Environment* No. 15, pp. 296 – 308.
- [29] Popa-Lala I., Anis C. N. (2010) The Assessment of the Company Financial Performances. In: *Proc. the 5th WSEAS International Conference on Economy and Management Transformation (Volume II)*, Wisconsin, USA, WSEAS, pp. 756-761.
- [30] Renewed EU strategy 2011-14 for CSR (2011) on EC Enterprise and Industry. [Online]. Available: http://ec.europa.eu/enterprise/policies/sustainable-business/files/csr/new-csr/act_en.pdf.
- [31] Schaltegger S., Wagner M. (2006) Integrative Management of Sustainability Performance, Measurement and Reporting, *International Journal of Accounting, Auditing and Performance Evaluation*. Vol. 3, pp. 1-19.
- [32] Stávková J., Souček M., Stojarová Š. (2009) Rela system. Software, Mendel University, Brno, Czech Republic.
- [33] Stefan V., Duica M., Coman M., Radu V. (2010) Enterprise Performance Management with Business Intelligence Solution,” in *Proc. 4th WSEAS International Conference on BUSINESS ADMINISTRATION (ICBA '10)*, Wisconsin, USA, WSEAS, pp. 244-250.
- [34] Sustainability Reporting Framework website (2012) on Globalreporting. [Online]. <https://www.globalreporting.org/reporting/reporting-framework-overview/Pages/default.aspx>.
- [35] Taticchi P., Tonelli F., Cagnazzo L. (2009) Development of a Performance Measurement System: Case Study of an Italian SME, In: *Proc. E-ACTIVITIES'09/ISP'09, 8th WSEAS International Conference on E-Activities and information security and privacy*, Wisconsin, USA, WSEAS, pp. 42-47.
- [36] Towards Integrated Reporting - Communicating Value in the 21st Century (2011) on IIRC. [Online]. Available: http://theiirc.org/wp-content/uploads/2011/09/IR-Discussion-Paper-2011_spreads.pdf.
- [37] The United Nations Environment Programme Finance Initiative website. (2012) [Online]. Available: <http://www.unepfi.org/>.
- [38] The United Nations Environment Programme Sustainable Buildings and Climate Initiative website. (2012) [Online]. Available: <http://www.unep.org/sbci/>.
- [39] The United Nations Global Compact website. (2012) [Online]. Available: <http://www.unglobalcompact.org/>.
- [40] The US G.A.A.P website (2012). [Online]. Available: <http://cpaclass.com/gaap/gaap-us-01a.htm>.
- [41] Varian H. R. (1992) *Microeconomic Analysis*, 3rd ed., New York: W. W. Norton and Company.
- [42] The XBRL website (2012). [Online]. Available: <http://www.xbrl.org/>.
- [43] Marty M., Socially Responsible Investing: United Nations Principles. [Online]. Available: http://works.bepress.com/cgi/viewcontent.cgi?article=1010&context=marty_martin.
- [44] CSR website. Available: http://ec.europa.eu/enterprise/policies/sustainable-business/corporate-social-responsibility/index_en.htm.
- [45] Czech Statistical Office. [Online]. Available: http://vdb.czso.cz/vdbvo/en/maklist.jsp?kapitola_id=33&expand=1.
- [46] Kocmanová, A., Dočekalová, M. (2012) Construction of the economic indicators of performance in relation to environmental, social and corporate governance (ESG) factors. *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis*, Vol. 60, No. 4, pp. 203-209.
- [47] Hodinka, M., Štencl, M., Hřebíček, J., Trenz, O. (2012) Current trends of corporate performance reporting tools and methodology design of multifactor measurement of company overall performance. *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis*, Vol. 60, No. 2, 2012, pp. 85-90.
- [48] Reference Document on Best Environmental Management Practice in the building and construction sector. Final Report, September 2012. JRC. Available <http://susproc.jrc.ec.europa.eu/activities/emas/documents/ConstructionSector.pdf>.
- [49] Green Paper. [Online]. Available: <http://eur-lex.europa.eu/LexUriServ/LexUriServ.do?uri=COM:2011:0164:FIN:EN:PDF>.

- [50] Investor Demand for Environmental, Social and Governance Disclosures, IFAC (2012) [Online]. Available: <http://www.ifac.org/publications-resources/investordemand-environmental-social-and-governancedisclosures>.
- [51] Chvátalová Z., Šimberová I. (2012) Analysis and identification of joint performance measurement indicators: social and corporate governance issues. *Acta Universitatis Agriculturae et Silviculturae Mendelianae Brunensis*, LX, No. 7, pp. 127–138.
- [52] Guidance on Corporate Responsibility Indicators in annual Reports (2012) [Online]. Available: http://unctad.org/en/Docs/iteteb20076_en.pdf.

Maple - pro marketing a inovace - nástroj měření výkonnosti podniku

Zuzana Chvátalová, Pavlína Charvátová

Vysoké učení technické v Brně, Fakulta podnikatelská

Kolejní 4, 612 00 Brno

chvatalova@fbm.vutbr.cz, xpcharv02@std.fbm.vutbr.cz

Abstrakt

Příspěvek diskutuje potřebu užití kvantitativních metod v paralele s využitím vhodného software pro stanovení / modifikaci strategie podniku a zvyšování jeho výkonnosti. Jsou zmíněny některé aktuální vlastnosti systému Maple 16. V krátké studii jsou jako příklad prezentovány vybrané skutečnosti, které se řeší v závěrečné práci obhájené na Fakultě podnikatelské Vysokého učení technického v Brně spoluautorky tohoto příspěvku. Jde o ukázkou implementace nástrojů systému Maple 16 při hodnocení marketingového výzkumu podniku působícího nedlouhou dobu na trhu s ohledem na možnost aplikace měřeného postupu v budoucnosti.

Abstract

The paper discusses the need for the use of quantitative methods in parallel using the appropriate software for setting / modification of enterprise strategy and increasing its performance. The article refers some current attributes of Maple 16. In a short study, we present as an example the chosen facts that co-author of this paper addressed in the final thesis defended at the Faculty of Business and Management of the Brno University of Technology. This is a sample the implementation of Maple tools for evaluation of enterprise marketing research, which operates only short time on the market. It is with respect to the possible application of the measurement process in the future.

Klíčová slova

Maple, marketing, regrese, statistika, strategie

Keywords

Maple, marketing, regression, statistics, strategy

1. Úvod

„Mít úspěch znamená mít vizi, mít myšlení zítřka, vidět do budoucnosti.“ Jan Moulis (Capital Partners)

² Implementovat prostředky informačních a komunikačních technologií (ICT), informačních databází a sociálních sítí podporuje rozvíjení schopností a dovedností nejen top manažerů, ale i zaměstnanců na všech úrovních podniku a podporuje jejich potřebu širší vzdělanosti, resp. rekvalifikace. To vede ve svém důsledku ke zdokonalování procedur v podniku, zefektivnění procesů a časovým úsporám. Ale i k propojování vztahů uvnitř i vně podniku. Dále k podpoře inovací, realizací a vyhodnocování průzkumů, zvláště marketingových výzkumů. Informační věk vybízí ke kvantifikaci (kvantifikovatelných, ale i kvalitativních faktorů a znaků), a tím k vývoji systémů nástrojů pro možnou měřitelnost a srovnání (podniků, výkonů, ekonomik aj.). Tato fakta ucelují formování strategie podniku. K tomu je však nezbytná investice do rozvoje lidského faktoru a do implementace technologického pokroku.

Tedy zužitkovat i neviditelná aktiva je výrazně významné pro rozhodování, více v (6).

Pokročilý model pro měření výkonnosti firmy s ohledem na strategický význam vyvíjel ve svém projektu David Norton s konzultantem Robertem Kaplanem (konec roku 1990). Výsledky projektu jsou shrnuty v článku „Balanced Scorecard – Measures That Drive Performance“ publikovaném v Harvard Business Review (leden – únor 1992). Důležitost výběru správných měřítel

² euroekonom.cz, ekonomický portál. *Svět ekonomie, obchodu a politiky v citátech* [online]. [cit. 2012-09-01].

Dostupný z WWW: < <http://www.euroekonom.cz/citaty.php> >

(s ohledem k naplnění strategie) byla následně popsána v článku „*Putting the Balanced Scorecard to Work*“ publikovaném rovněž v Harvard Business Review (září – říjen 1993), kde byly zhodnoceny další Nortonovy praktické zkušenosti a aplikace. Tyto zkušenosti ukázaly, že již přibližně dvacet až dvacet pět měřítek napříč čtyřmi perspektivami (finanční; zákaznické; interní; inovační a růstové) podniku umožňují formulovat a implementovat svou strategii. Komplexní přístup k rozhodování umožní zjistit i příčinné souvislosti, které charakterizují strategickou trajektorii – především význam investice do zvyšování kvalifikace zaměstnanců, do implementace informačních a komunikačních technologií, inovací, více v (7).

Na konci minulého století rostoucí automatizace a požadavky na vyšší produktivitu snížily počet lidí vykonávající tradiční pracovní úkony, zatímco rostoucí konkurenční tlak vyžaduje více pracovníků specializovaných na analytické funkce: *engineering, marketing, management a administrativu*. Kladně jsou hodnoceny podněty jedinců i z výrobní sféry. Podniky se snaží transformovat tak, aby byly konkurenceschopné i v budoucnosti. Manažeři potřebují k řízení podniku určitý soubor nástrojů k vyhodnocování činnosti podniku, jeho ekonomického prostředí a dosahování svých cílů. Takový soubor nástrojů poskytuje například tzv. *Balanced Scorecard* (BSC, tzv. vyvážený scorecard), který manažeři mohou využít pro navigaci budoucího úspěchu, a tím pro stanovení cílů a metod k jeho dosažení. BSC umožňuje nejen sledovat finanční výsledky ale i schopnost podniku jistit *nehmotná aktiva* k budoucímu růstu. Konkurenční prostředí informačního věku se mění a vyostřuje. Informační věk přináší nové podmínky pro konkurenci. (7)

Nová provozní prostředí podniků v éře informačního věku jsou determinována novými provozními podmínkami, tzv. *křížovými funkcemi*, spojením zákazníky s dodavateli, segmentací trhu, globalizací, inovacemi, znalostními pracovníky, více v (7).

„*Stroje jsou od toho, aby běžely automaticky. Úkolem lidí je myslet, řešit problémy a zajišťovat jakost, ne sledovat součástky na pásu. V našem případě jsou lidé řešiteli problémů, nikoli variabilními náklady.*“³ (8)

Dle Petera Druckera „*Podnikání má dvě - a pouze dvě - základní funkce: marketing a inovaci. Marketing a inovace plodí výsledky, vše ostatní jsou náklady.*“ (1)

Řada zásadních i podpůrných firemních rozhodnutí podléhá zkušenosti, odborné úrovni firemního managementu, a tedy jeho podporou pro zvyšování odborné gramotnosti zaměstnanců. Dle Benjamina Franklina „*Géníus bez vzdělání je jako stříbro v dole.*“ (1) Pouze intuice a risk či investice do služby externího experta nemusejí přinést kýžený úspěch. Na druhé straně rozhodnutí mají stále častěji interdisciplinární charakter vyžadující erudované kroky. Jde například o využití pokročilých kvantitativních metod, metod soft computingu, citlivé dolování dat i korektně provedené vědecké výpočty. (2) Tím více, že v důsledku „boomu“ nástrojů ICT jsou v současnosti kvantitativní metody úspěšné i v oborech, které dříve byly považovány výhradně za společensko-vědní. (5)

Tedy implementace vhodných ICT přispívá k tomu, aby erudovaná rozhodnutí byl schopen učinit uživatel s patřičnou úrovní znalostí.

Bill Gates řekl: „*Lepší, než předpovídat budoucnost, je vytvořit ji,*“ (1) Fakulta podnikatelská Vysokého učení technického v Brně (FP VUT) z pohledu technologických inovací patří k velmi dobře a moderně vybaveným vysokoškolským pracovištím. Jde nejen o hmotné technické vybavení, ale i o software. Od svého vzniku fakulta preferuje a cíleně pěstuje svou dobrou pověst na bázi odrazu kvality svých aktivit v propojení teoretického a výzkumného potenciálu s praxí. Jde především o reálné aplikace získaných vědomostí nejen v okamžiku aktivního studia posluchačů, avšak v téže míře i o možnost propojení s jejich budoucností. Již řadu let jena FP VUT možné zpracovávat problematiku v prostředí počítačového systému Maple, kterým jsou vybavena vybraná pracoviště fakulty, posluchárny a laboratoře. Systém Maple je zaměřen jak na výuku, tak na využití při komplexním řešení celé řady skutečných problémů nejrůznějších oborů. Stejně tak jsou podporováni i akademičtí pracovníci v rámci svých pedagogických, výzkumných i vědeckých aktivit. I tím bezesporu fakulta má možnost sekundárně působit, resp. ovlivňovat i dění v samotném podniku, s nímž posluchač či akademik spolupracuje. Níže bude uvedena ukázka vybraných řešených aspektů podniku krátce

³ Současné připomeňme myšlenku Henryho Forda: „*Myšlení je nejtěžší práce, jaká existuje. To je pravděpodobně důvod, proč tak málo lidí myslí.*“ (1)

působícího na trhu, kterými mj. se zabývala práce (3). Přičemž přínos šetření nespočívá jen v samotné okamžité analýze situace. S podporou implementace počítačového systému jde především o návrh procedur (v jistém slova smyslu o algoritmizaci) vyhodnocování situace pro nové vstupy získané v budoucích měřeních a průzkumech. Jejich aplikace je téměř automatická. Výstupy tak poskytnou managementu aktualizované, přitom rozumně dostupné informace o stavu firmy. Management pak získá solidní podklady pro erudované rozhodování o případné změně a vedení firemní strategie v budoucnosti.

2. Maple 16

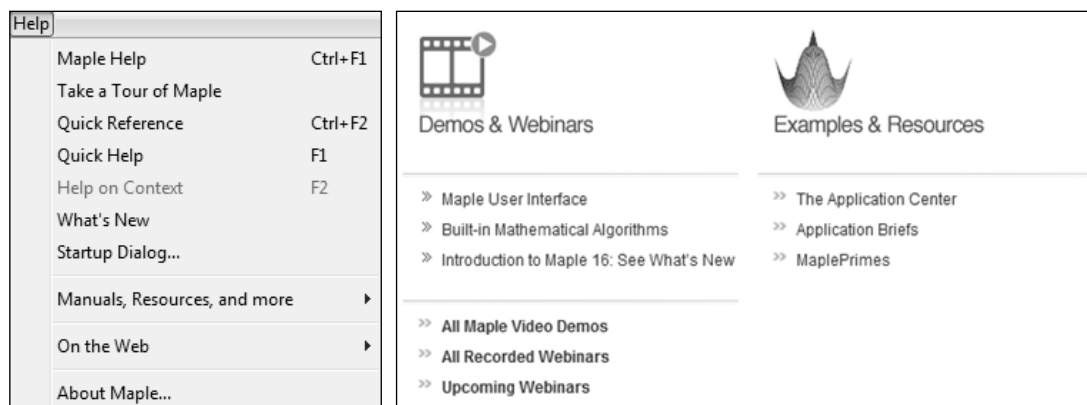
O počítačovém systému Maple bylo řečeno na nejrůznějších fórech již mnoho. Kompletní a aktualizované informace jsou dostupné na webových stránkách <http://www.maplesoft.com> společnosti Maplesoft, Inc., která tento systém vyvíjí. Oblíbenost nasazení systému prostupuje v celosvětovém měřítku nejrůznějšími skupinami uživatelů. (4) Zmíňme krátce jeho vybrané vlastnosti s ohledem na současnou verzi Maple 16. Vývoj systému interaktivně reaguje na požadavky svých uživatelů a společenského pokroku⁴.

Aktuální verze *Maple 16* byla expedována v roce 2012. Obsahuje více než čtyři tisíce pět set doplňků a zdokonalení klíčových oblastí systému. Jsou to zejména: vylepšení uživatelského rozhraní, tzv. klikací matematika (*Clickable Math*™ 3.0), výpočetní účinnost, zrychlení jádra Maple, paralelní výpočty a matematické algoritmy, programovací jazyk, vysoce kompaktní vizualizace, „chytré“ 2-D grafické prohlížeče, pružné zvětšování, správa paměti. Dále jde o nástroje pro výuku i pro využití v praxi, konkrétně o více než sto vylepšení matematických aplikací (*Math Apps*) především v oblastech algebry a geometrie, diferenciálního a integrálního počtu, funkcí a relací, kreslení grafů, trigonometrie, pravděpodobnosti a statistiky, v oborech fyziky, financí a ekonomie.

Pracovní prostředí Maple je uživatelsky snadno obslužné. Zápisy výrazů a struktur lze vést užíváním obsáhlých kontextových palet nástrojů intuitivně v analogii s pravidly standardních zápisů rukou. Pomocí interaktivních asistentů a nástrojů lze obrázky, výpočty i komentáře doplňovat a upravovat, dokumenty strukturovat, zabudované aplikace modifikovat a vkládat. K vkládání textu do obrázku lze použít i přímý příkaz `textplot(L, options)`. Výpočty i grafy lze „oživovat“ pomocí posuvníků s možností ukládání příslušné syntaxe, zdrojového kódu, a to jednoduchou aktivací v nabídce v horní liště Maple zápisníku.

Nápověda a navigace v Maple je srozumitelná, rozsáhlá a nabízená z několika pohledů – od elementárních kroků zabudovaného Help podsystemu, rychlých navigací, včetně implementovaných knihoven s teoretickým zázemím a celou řadou příkladů, až po demoverze videí a odborníky vedené pravidelné diskuze a webináře (Obr. 1). Tedy systém Maple je dostupný celému spektru uživatelů od naprostých začátečníků až po extrémně pokročilé uživatele využívající nejsložitějších nástrojů tohoto systému jak v praxi, tak i ve vědě.

⁴ Základní zdroj celého tohoto odstavce: Mpalesoft. [online] 2012, [cit. 2012-1-10]. Dostupný z: <http://www.maplesoft.com/solutions/engineering/>

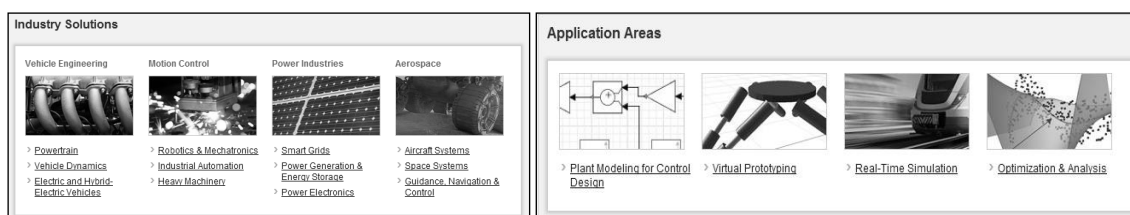


Obr. 1: Vybrané prostředky nápovědy v Maple

[Zdroj: Zabudovaný manuál v systému Maple > Help (vlevo); Mpalesoft. [online] 2012, [cit. 2012-1-10]. Dostupný z: <http://www.maplesoft.com/products/Maple/> (vpravo)]

Podpora využití Maple je stále více konkretizována, rozšiřována a zdokonalována při zohlednění základního členění uživatelů do skupin: *studenti*, *akademici* a *profesionálové*. Je zrychlena analytická i aplikační pracovní produktivita, zdokonaleno výkonné jádro Maple, šíře a hloubka prezentace výstupů samotného systému. Dále je podporována možnost sofistikovaných návrhů simulací a modelování (*MapleSim 6*). Expanduje množství zabudovaných a předdefinovaných „užitečných“ funkcí a procedur cílených většímu počtu odborných profesí a pracovních zařazení v těchto profesích.

Maple 16 pro profesionály (Maple 16 for Professionals) rozvíjí především následující oblasti: elektroniku, fyziku, energetiku, vzdušný a kosmický prostor, automobilový průmysl, finanční modelování a ekonomii, statistiku a procesní řízení, velmi výkonné výpočty, operační výzkum, zpracování signálu, virtuální vytváření prototypů, simulace v reálném čase, optimalizace, řízení pohybu aj. (Jmenujme několik vybraných uživatelů systému Maple: Toyota Motor Corp., Renault, Ulysse Nardin, Magtech AS, Arqiva, Marquardt GmbH, Aerospace Manufacturing Technology Centre (AMTC).) Některé oblasti aplikací a řešených problematik v systému Maple uvádí Obr. 2.



Obr. 2: Vybrané ukázky řešených problémů a aplikačních oblastí v Maple

[Zdroj: Mpalesoft. [online] 2012, [cit. 2012-1-10]. Dostupný z: <http://www.maplesoft.com/solutions/engineering/>]

Maple 16 pro akademiky (Maple 16 for Academic) je nástrojem pro analýzy, výzkumy, vizualizace a řešení matematických problémů (pro nejrůznější obory matematiky). Intuitivní rozhraní podporuje více způsobů interakce. Je využíván programovací jazyk, podobný jazyku Pascal. Pomocí tzv. „chytrého“ dokumentu (*Smart Document*) lze automaticky dokumentovat posloupnost výpočtů, podpůrné použité technické znalosti, vysvětlující text, grafy, obrázky, zvuky, fotografie, diagramy i poznámky v elektronické podobě. Zdroje pro práci akademiků podporují následující centra: *The Teacher Resource Center* (<http://www.maplesoft.com/TeacherResource/index.aspx>), *The Application*

Center (<http://www.maplesoft.com/applications/index.aspx>), *The Teaching Concepts with Maple The Teaching Calculus with Maple: A Complete Kit, Application Briefs, MaplePrimes* (<http://www.mapleprimes.com/>; <http://www.maplesoft.com/applicationbriefs/>). Kromě výše, resp. i níže uvedených vylepšení Maple 16 nabízí akademikům zejména zdokonalené vizualizace (2D a 3D grafy a animace), editor rovnic, interaktivní asistenty, tutorý a kontextová menu, „chytré“ palety (*Smart Popups*), šablony pro řešení úloh, nástroj „uchop-řeš“ (*Drag-to-Solve*). Vyšší kvalitu nabízí v Maple 16 i variabilní manažer, paměťový správce a průzkumný asistent, dokument pro práci s posluchači tzv. *MapleCloud*, rozpoznávač ručně psaných symbolů, zpracování „živých“ dat, konektivita, balíček paralelního programování aj. Do Maple dokumentu lze snáze i vkládat interaktivní komponenty (posuvníky, tlačítka, voliče stupnic, měřidla, technické zprávy a další). Nové aplikace verze Maple 16 akademici ocení především při výuce v předmětech matematického a inženýrského základu na vysokých školách a nástavbách včetně celoživotního vzdělávání (mj. zmiňme zjednodušení testování a hodnocení posluchačů, statistiku výsledků vědomostí studentů), a v podpoře aplikovaného výzkumu a mnoha vědeckých aktivit. Maple respektuje obecný technologický pokrok a potřeby informační generace, jak ukazuje Obr. 3.



Obr. 3: Maple přehrávač pro iPad

[Zdroj: (Mpaesoft. [online] 2012, [cit. 2012-1-10]. Dostupný z:
<http://www.maplesoft.com/products/mapleplayer/>]

Maple 16 pro studenty (Maple 16 for Students) nabízí řadu nových vlastností a vylepšení zejména v oblasti symbolické a numerické matematiky v reálném a komplexním oboru, řešičů rovnic (i diferenciálních), teorie funkcí (zejména limit a polynomů), lineární algebry, diferenciálního a integrálního počtu, vektorového počtu, optimalizace, operační analýzy, programování, převádění jednotek a rozměrů, statistiky a odchylek. Velké množství studentských aplikací s volným přístupem je soustředěno v již výše zmíněném Aplikačním centru a pomocných podpor studentům přímo od expertů pro Maple systém v diskuzním fóru MaplePrimes. Maple 16 poskytuje nové studentské video záznamy, tutoriály, elektronické knihy, příručky a podpory (např. *Mathematics Survival Kit*) v oblasti základní, vyšší i pokročilé inženýrské matematiky, *Student Help Center* (<http://www.maplesoft.com/studentcenter/index.aspx>) a jako paralelní produkt *Maple 16 Portal* - portál pro studenty (Obr. 4), aj.



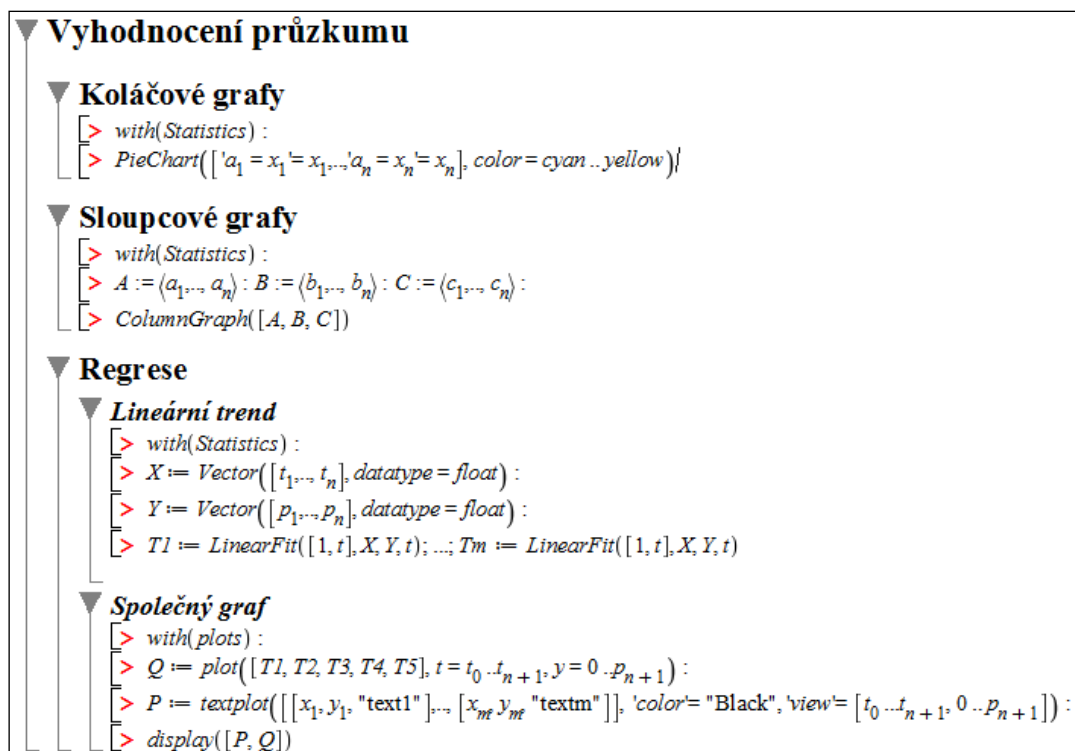
Obr. 4: Záhloví Maple 16 Portal

[Zdroj: Systém Maple™ Portal]

3. Maple: nástroj pro současné i budoucí využití při volbě strategie podniku – ukázka případové studie

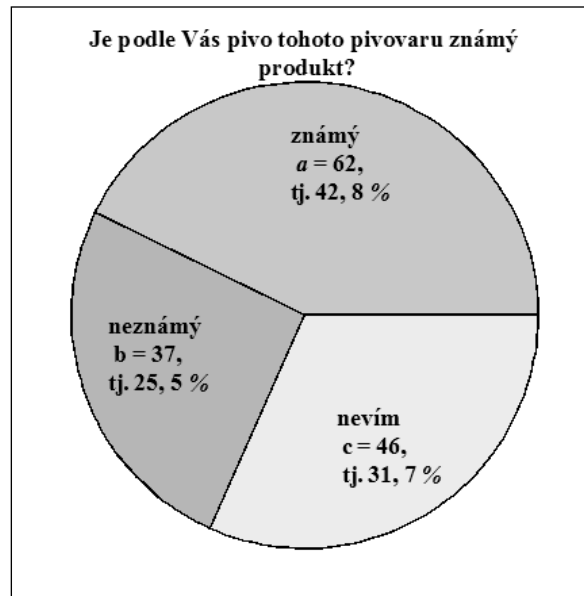
V tomto odstavci jsou v souladu s (3) uvedeny vybrané ukázky realizace vyhodnocení marketingového průzkumu s využitím Maple pro potřeby managementu malého podniku. Jde o zcela nový, moderní pivovar působící na trhu přibližně tři roky. Činnost pivovaru kloubí moderní techniku s tradiční technologií. Zaměstnává šestnáct zaměstnanců (ředitel, sládek, účetní, dva obchodní zástupci, servisní technik, údržbář, šest dělníků, marketing manažer, expedientka, řidič). Analýzami vnějšího i vnitřního okolí pivovaru byla identifikována řada zjištění, například: konkurenční výhodou je výborná chuť piva nepodporovaná konzervanty ani oxidanty a způsob dozrávání piva (v ležáckých tancích). Byl proveden *marketingový mix* (produkt, cena, distribuce, propagace), dále monitorování *dodavatelé, odběratelé, konkurence*. Z následné *SWOT analýzy* vyplývají: *silné stránky* (vysoká kvalita produktu, výborná chuť piva, lokalizace pivovaru, poutavé firemní barvy), *slabé stránky* (nedostatečná propagace, neznámý produkt, chybí pivovarská restaurace, krátká trvanlivost piva), *příležitosti* (zvýšení povědomí o společnosti, zvýšení počtu zákazníků, rozšíření pivovaru, vaření nového druhu piva, bio pivo, vyvážení piva do zahraničí), *hrozby* (vznik nového pivovaru v regionu, silná konkurence, změna legislativy, růst cen pohonných hmot, nespokojenost zákazníků). (3) Pro volbu firemní strategie, možnost inovací a činnost marketingového oddělení s cílem přispívat k postupnému zvyšování výkonnosti firmy vyvstala potřeba kvantifikovat situaci, v níž se nachází pivovar. Pro orientaci, resp. vysledování závislostí a časových vývojů příslušných veličin pak je třeba získaná fakta vhodně, přitom nekomplikovaně vizualizovat. I přes poměrně krátkou dobu působení pivovaru na trhu bylo rozhodnuto provést empirické šetření (byť s málo početnou časovou řadou údajů). Management vyžaduje marketingový průzkum v terénu orientovaný směrem ke spokojenosti zákazníků (s produkty, kvalitou, cenou a dalšími atributy s ohledem na signály výše uvedených kvalitativních zjištění). Je žádán průzkum podporovaný kvalitním teoretickým zázemím, dobře zpracovaný, dávající srozumitelné výstupy pro korektní ekonomickou interpretaci a možnost vyslovení nápravných doporučení při modifikaci budoucí strategie pivovaru managementem. Faktickým cílem celé akce (v současnosti i vzhledem k nízké početnosti údajů v časových řadách pozorovaných znaků) kromě zhodnocení současného stavu je návrh a standardizace postupných kroků (algoritmizace) vyhodnocujících průzkum s využitím vhodného počítačového prostředí a kvantitativních metod tak, aby byla možná znovurealizace obdobného průzkumu v budoucnosti. Přitom, aby bylo možné zpracování provést již v interních podmínkách podniku.

Dotazníkového šetření se zúčastnilo téměř sto padesát respondentů. Výhodou systému Maple je, že je kompatibilní s knihovnickými aplikacemi tabulkového procesoru Microsoft Excel, kam byla primární data získaná v terénu uložena. Data tak byla následně snadno načtena do Maple zápisníku. Dotazník obsahoval dvacet šest otázek různého typu. Při sumarizaci byly uvažovány jak absolutní četnosti, tak relativní četnosti odpovědí. V Maple je možné dokument přehledně strukturovat a vytvářet podsekce, které je možno sbalit nebo dále hypertextově rozbalit do nižších úrovní. Toho bylo využito při návrhu algoritmizace. Přitom do obecných příkazů na Obr. 5 stačí mechanicky doplnit empiricky zjištěné hodnoty a pak již nechat systém Maple samostatně pracovat. Z důvodu rozsahu příspěvku Obr. 5 zachycuje pouze vybrané sekce, podsekce a postupy.

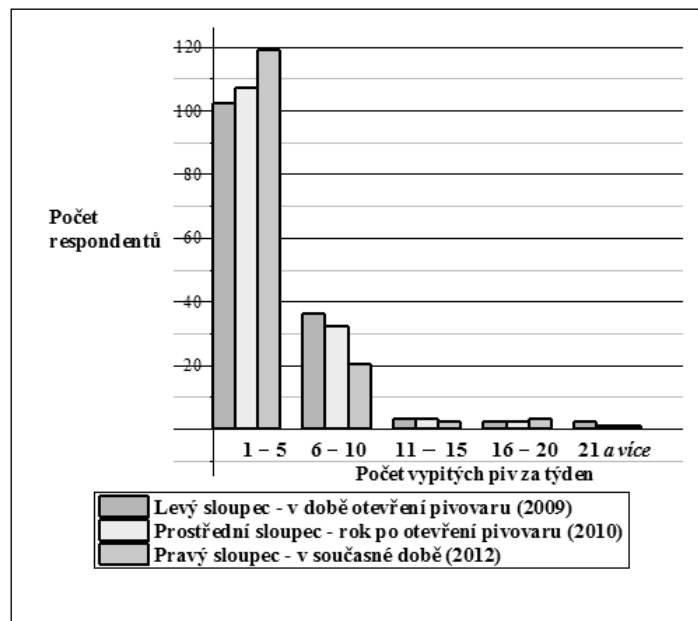


Obr. 5: Strukturovaný Maple zápisník do sekcí a podsekcí včetně příkazů pro algoritmizaci
[Zdroj: Vlastní zpracování v Maple]

- Otázky typu: *Je podle vás pivo z tohoto pivovaru známý produkt?* s předem danými variantami odpovědí (v tomto případě) „ano“ (toto pivo je známý produkt), „ne“ (myslím, že toto pivo lidé neznají) a „nevím“ (vůbec nemám tušení, zda je toto pivo známé) byly vizualizovány koláčovými grafy (*PieChart*), Obr. 6. Lze srozumitelně provést jejich vizuální vyhodnocení a s výstupy budoucích průzkumů snadno porovnat (výseče zachycující absolutní četnosti odpovědí současně odpovídají i relativním četnostem). Maple disponuje velmi pohodlnými interaktivními prostředky pro případnou modifikaci vlastností koláčů (popisků, barev, stylů aj.) užitím pravého tlačítka myši či kontextového menu v horní liště dokumentu.
- Otázka typu: *Kolik pŕllitrů piva z tohoto pivovaru v průměru za týden vypijete?* s předem stanovenými kategoriemi počtů vypitých pŕllitrů piva (1 až 5, 6 až 10, 11 až 15, 16 až 20, 21 a více). Navíc odpovědi jsou sledované ve třech časových momentech, a to v době otevření pivovaru (rok 2009), po jednom roce působení pivovaru na trhu a v současné době (rok 2012). Zjištěná data byla v Maple zapsána do kontingenční tabulky. Její třidimenzionální charakter byl převeden do dvoudimenzionálního typu užitím přehledného „vícenásobného“ sloupcového grafu (*ColumnGraph*), Obr. 7. Obdobně jako v předchozím případě je tento graf v Maple snadno modifikovatelný užitím interaktivních komponent systému.



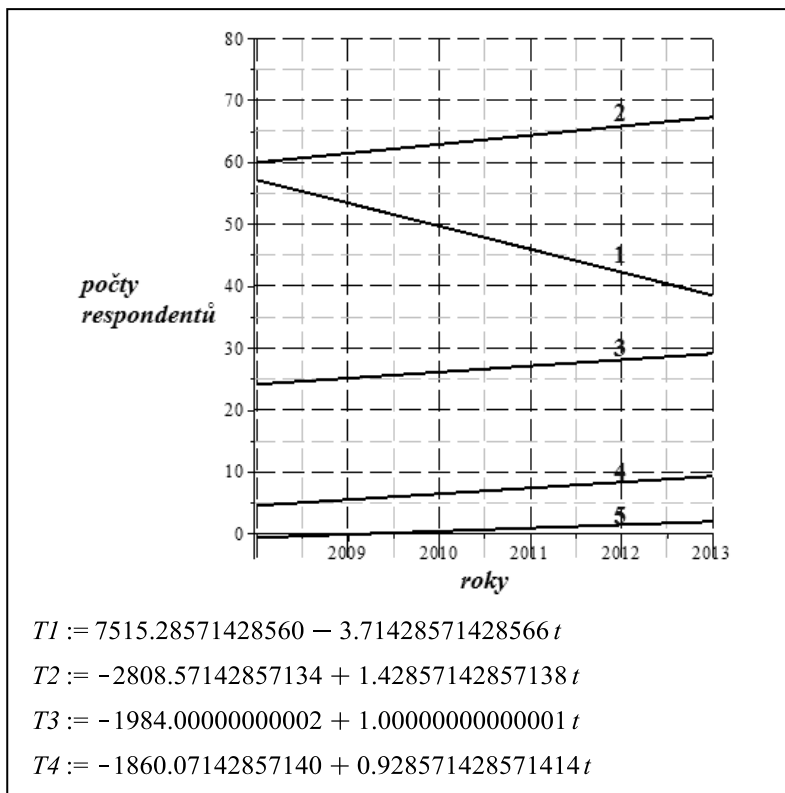
Obr. 6: Koláčový graf v Maple
[Zdroj: (3) a vlastní zpracování v Maple]



br. 7: Slupcový graf v Maple
[Zdroj: (3) a vlastní zpracování v Maple]

- Nyní se zaměříme na jinou možnost vyhodnocení předchozího typu otázky. Budeme sledovat trend vývoje jednotlivých variant odpovědí v čase. K tomu využijeme v Maple zabudované prostředky statistického balíčku (*Statistics*) a balíčku pro kreslení grafů (*plots*). Zhodnotíme otázku: *Ohodnoťte pivo tohoto pivovaru jako ve škole známkou 1 až 5*. Rovněž šlo o hodnocení i v čase, a to v době otevření pivovaru (rok 2009), rok po otevření pivovaru (2010) a v současné době (rok 2012). Pro zjištění trendu byla zvolena lineární regrese. Vývoj všech uvažovaných variant odpovědí byl pro zprůhlednění, snazší porovnání a analýzy zachycen do společného souřadnicového systému (Obr. 8). Předpisy pro příslušné lineární regresní přímky (vypočítané v systému Maple) využitím příkazu *LinearFit* ze statistického balíčku jsou rovněž

vyjádřeny v Obr. 8. Indexace je přitom v korespondenci s přiděleným „oznámkováním“ a totéž respektuje označení přímek ve společném grafu. Je vidět, že společný graf je velmi srozumitelný pro interpretaci empirických zjištění. V souladu s výše uvedenými analýzami (SWOT aj.) pak tato kvantifikovaná zjištění hrají významnou roli pro výstavbu firemní strategie. Poznamenejme, že v systému Maple je možno dále provést statistickou diagnostiku získaných modelů. A pro další, resp. složitější statistická šetření nad získanými daty využít celé řady zabudovaných statistických nástrojů, například pro zjišťování závislostí znaků či testování hypotéz. To z důvodu rozsahu příspěvku toto neuvádíme.



Obr. 8: Regresní lineární modely v Maple
[Zdroj: Vlastní zpracování v Maple]

4. Závěr

Manažer podniku je často nucen učinit rychlá a interdisciplinární rozhodnutí. Má-li v ruce nástroj, který umí solidně ovládat, přitom sám disponuje vědomostmi tak, aby mohl odborně korektně komunikovat s jeho rozhraním a výstupy pak odpovědně interpretovat, stává se ve svých rozhodnutích více nezávislým. Takový nástroj spolu s vědomostmi, zkušenostmi a intuicí manažera lze považovat za konkurenční výhodu firmy. Systém Maple spolu s aplikací kvantitativních metod takovou příležitost poskytuje. Je však potřebné, aby již v etapě vzdělávání studenti získávali správnou počítačovou gramotnost a budovali v sobě potřebu se kontinuálně vzdělávat, rekvalifikovat apod. Fakulta podnikatelská vysokého učení technické v Brně je dobrým příkladem takových snah.

5. Poděkování

Příspěvek vznikl s přispěním projektu „Konstrukce metod pro vícefaktorové měření komplexní podnikové výkonnosti ve vybraném odvětví“ GAP403/11/2085 podporovaného Grantovou agenturou České republiky.

6. Literatura

- [1] euroekonom.cz, ekonomický portál. *Svět ekonomie, obchodu a politiky v citátech* [online]. [cit. 2012-09-01]. Dostupný z WWW: < <http://www.euroekonom.cz/citaty.php> >
- [2] Hřebíček J. a kol. Vědecké výpočty v matematické biologii. Brno: Akademické nakladatelství CERM, s. r. o. Brno, 2012.
- [3] Charvátová P. Měření spokojenosti zákazníků společnosti Pivovar Chotěboř s.r.o. užitím Maple. Brno: Vysoké učení technické v Brně, Fakulta podnikatelská, 2012. 110 s. Vedoucí diplomové práce RNDr. Zuzana Chvátalová, Ph.D.
- [4] Chvátalová, Hřebíček J. and Žigárdy M. Computer simulation of stock exchange behavior in Maple, International journal of mathematical models and methods in applied sciences, Vol.5, No.1, 2011, pp. 59-66.
- [5] Chvátalová Z., Hřebíček, J. Modeling and Simulation Utility Functions with Maple. Mendel Journal series. Vol. 18, no. 1, 2012. pp. 552-557.
- [6] Itami H. Mobilizing Invisible Assets. Cambridge, Mass.: Harvard University Press, 1987.
- [7] Kaplan R.S., Norton D. P. Balanced Scorecard. Strategický systém měření výkonnosti podniku. Praha: Management Presss, 2000.
- [8] Kaplan R.S., Sweeney A. Romeo Engine Plant. Boston: Harvard Business School, 1994.

Evoluční model pesimizace prostředí v Maple

Jiří Kalina

Masarykova univerzita, Institut biostatistiky a analýz,
Kamenice 126/3, Brno, 625 00
kalina@mail.muni.cz

Abstrakt

*Příspěvek se zabývá popisem a diskuzí populačního modelu se zahrnutím náhodné evoluční složky, založeného na jednoduchém modelu chemostatu popsáném v [2]. Model je rozšířen na neomezený počet populací a je z něj vyňata omezující podmínka rovnosti mortality a podílu nevyužitého substrátu podléhajícího zkáze, což umožňuje jeho aplikaci i na složitější společenstva. Jádrem příspěvku je implementace modelu v prostředí Maple, ve kterém je pomocí numerického řešiče diferenciálních rovnic *dsolve* spočten vývoj velikosti populací a na výsledných grafech je demonstrována oprávněnost předem vyslovených předpokladů.*

Abstract

This paper deals with the description and discussion of the population model including a random evolutionary component, based on a simple model of chemostat described in [2]. The model is extended to an unlimited number of populations and a condition of equality of mortality and the proportion of unused perishable substrate, which allows its application to more complex systems. The core contribution is an implementation of the model in a Maple environment with an use of a numerical differential equation solver "dsolve". The evolution of population sizes is counted and plots of results demonstrate the legitimacy of former assumptions.

Klíčová slova

Populační model, pesimizace.

Keywords

Population model, pesimization.

1. Základní vlastnosti modelu

Předpokládejme uzavřený systém složený z prostředí a n navzájem přímo neinteragujících populací, které ho obývají, $n \in N_0$. Jedince náležející k jedné populaci budeme souhrnně nazývat druhem. Úživnost prostředí, tj. veličinu, která charakterizuje podmínky pro přežití a reprodukci jedinců, označme S a považujme ji zjednodušeně za množství potravy (substrátu, resp. energie), které je k dispozici společně pro všechny populace v prostředí. S budeme považovat za funkci času $S(t)$. Přírůstek substrátu v prostředí za jednotku času (růst, resp. transport energie do systému) považujme za konstantní a označme jej d .

Vzhledem k předpokladu, že populace v modelu spolu přímo neinteragují, omezuje se jejich vzájemné ovlivnění na soupeření o dostupné množství potravy, jde tedy o ryze kompetitivní vztah [1]. Vzhledem k jednoduchosti modelu budeme druhy odlišovat pouze schopností jejich zástupců spotřebovávat společnou potravu a takto přijatý substrát využít k vlastnímu přežití a reprodukci. Tento proces popíšeme dvojicí veličin q a e , kde q_i představuje množství substrátu přijaté jedním jedincem i -té populace za jednotku času. Tuto veličinu budeme dále považovat za funkci množství dostupného substrátu $q_i(S)$. Naproti tomu hodnota e_i označuje efektivitu přeměny substrátu na živé jedince i -tého druhu jako poměr spotřebovaného substrátu vůči počtu nově narozených jedinců. V první variantě modelu budeme hodnoty q_i a e_i považovat za vzájemně nezávislé. Dále budeme předpokládat konstantní úmrtnost m , stejnou pro všechny populace, jako podíl jedinců daného druhu uhynulých za jednotku času a nakonec zbytkový podíl nespotebovaného substrátu z , který podlehně zkáze, aniž by byl spotřebován některou z uvažovaných populací.

Na počátku modelovaného období budeme předpokládat jednu (nebo větší konstantní počet) populaci v prostředí, která spotřebovává dostupný substrát pro vlastní reprodukci. Snadno lze ukázat, že po určité době dojde k ustálení počtu jedinců populace (uvědomme si, že jde o silně zjednodušený idealizovaný model) na hodnotě odpovídající úživnosti prostředí, která se v dalším průběhu času již nebude měnit. Je nutno poznamenat, že pro jednoduchost model opomíjí veškeré komplikace, které s sebou přináší rozmanitost pohlavního rozmnožování jako je samo pohlaví, hledání partnera či nestejný počet potomků každého jedince. Budeme tedy nadále předpokládat, že dostatečným požadavkem pro další rozmnožování populace je přítomnost alespoň jednoho jedince daného druhu v systému, že potomkem jedince určitého druhu je (až na výjimky) jedinec téhož druhu a že křížení mezi jedinci různých druhů není možné.

Situace se nicméně začne komplikovat ve chvíli, kdy se v prostředí objeví zástupce nového druhu. V modelovém uzavřeném systému není možný průnik takového jedince zvenčí, proto jej budeme považovat za mutanta některé z již existujících populací, který se bude odlišovat parametry q_i a e_i , poruší tedy předpoklad z předcházejícího odstavce. Výskyt mutantů bude relativně řídký a bude se odehrávat pouze v určitých okamžicích. Takový jedinec se v našem modelu může, bude-li úspěšný ve vzájemné kompetici, stát zakladatelem populace nového druhu, v opačném případě může ovšem nová populace záhy vyhynout a uvolnit tak prostředky pro zástupce jiných druhů.

Časovou změnu množství substrátu dostupného v systému je na základě uvedených předpokladů možné vyjádřit následujícím vztahem:

$$\frac{\partial S(t)}{\partial t} = d - z \cdot S(t) - \sum_{i=1}^n q_i(S) \cdot P_i \quad (1)$$

kde $S(t)$ je dostupné množství substrátu v čase t , d je přírůstek substrátu za jednotku času, z je podíl nevyužitého substrátu za jednotku času, n je počet populací přítomných v systému, $q_i(S)$ je spotřeba substrátu pro jednoho jedince i -tého druhu za jednotku času a P_i je počet jedinců i -té populace.

Popišme nyní tvar závislosti $q_i(S)$ množství substrátu spotřebovaného jedincem i -té populace za jednotku času. Opět zjednodušeně předpokládejme, že dostupnost substrátu v prostředí je stejná pro všechny jedince všech populací a její vyhledávání proto nehraje v modelu žádnou roli. V případě, že bude množství dostupného substrátu nulové, $S = 0$, bude přirozeně rovněž $q_i = 0$. Jakmile začne v prostředí narůstat množství dostupné potravy, poroste přímo úměrně také množství zkonsumované jedinci všech populací. Vzhledem k fyziologickým možnostem organismů je však maximální množství potravy přijaté za jednotku času shora omezené, proto se bude s narůstajícím množstvím substrátu v prostředí křivka funkce q_i oddělovat od přímky přímé úměry, až plynule konverguje ke konstantní funkci. Uvedený předpoklad dokonale splňuje v biologii a příbuzných oborech hojně využívaná logistická křivka následujícího tvaru:

$$q_i(S) = \frac{a \cdot S}{1 + B_i \cdot a \cdot S} \quad (2)$$

kde parametr a určuje strmost počáteční části křivky a rovněž rychlost dosažení následující asymptotické fáze, zatímco parametr B_i její výšku. V první, jednodušší, variantě modelu budeme předpokládat jednak, že B_i je totožné pro všechny populace (což není sám o sobě předpoklad, který by výrazně ovlivnil chování modelu), současně ale také, že B_i je konstantní a nezávislé na ostatních parametrech modelu. Rovněž parametr a budeme v první variantě modelu považovat za konstantu.

Přejdeme nyní k rovnici popisující vývoj počtu jedinců i -té populace v čase. Změna počtu jedinců, P_i bude vyjádřena jako množství substrátu zkonsumovaného jedinci populace, vynásobeného efektivitou přeměny substrátu na živé jedince, minus počet uhynulých jedinců:

$$\frac{\partial P_i(t)}{\partial t} = q_i(S) \cdot e_i \cdot P_i(t) - m \cdot P_i(t) \quad (3)$$

kde $P_i(t)$ je počet jedinců populace v čase t , $q_i(S)$ je spotřeba substrátu pro jednoho jedince i -tého druhu za jednotku času, e_i je efektivita přeměny substrátu na živou hmotu i -tou populací a m je relativní úmrtnost, pro jednoduchost opět totožná pro všechny populace.

Uvedený postup nám poskytuje soustavu $n + 1$ diferenciálních rovnic, jejímž řešením lze získat informaci o vývoji počtu jedinců jednotlivých populací v čase. Vzhledem k tomu, že v první variantě

modelu uvažujeme q_i totožné pro všechny populace v systému, je řešení závislé především na hodnotách e_i pro jednotlivé populace. V případě, že se v systému nevyskytne mutant, lze ukázat, že v dlouhodobém horizontu se v závislosti na hodnotách konstantních parametrů d, z, a a b_i a na počátečních podmínkách $S(0)$ a $P_i(0)$ v systému ustálí stav, ve kterém je množství dostupného substrátu v prostředí konstantní a navíc počet jedinců nejvýše jedné populace nekonverguje k nule.

Platí totiž:

$$\frac{d\left(S + \sum_{i=1}^n \frac{P_i}{q_i}\right)}{dt} = d - m\left(S + \sum_{i=1}^n \frac{P_i}{q_i}\right) - (m - z) \sum_{i=1}^n \frac{P_i}{q_i} \quad (4)$$

a tedy po integraci podle t potom:

$$S + \sum_{i=1}^n \frac{P_i}{q_i} = \int e^{zt} d \, dt \cdot e^{-zt} + \int e^{zt} (z - m) \sum_{i=1}^n \frac{P_i}{q_i} dt \cdot e^{-zt} + C \cdot e^{-zt} \quad (5)$$

a po úpravě:

$$S + \sum_{i=1}^n \frac{P_i}{q_i} = \frac{d}{z} + (z - m) \int e^{zt} \sum_{i=1}^n \frac{P_i}{q_i} dt \cdot e^{-zt} + C \cdot e^{-zt} \quad (6)$$

za předpokladu, že pro $t \rightarrow \infty$ konvergují 2. a 3. sčítanec k nule, dostáváme výsledek:

$$\lim_{t \rightarrow \infty} \left(S + \sum_{i=1}^n \frac{P_i}{q_i} \right) = \frac{d}{z} \quad (7)$$

Tedy, pokud bude úživnost prostředí dostatečná, zvítězí v kompetici o omezený zdroj potravy v uzavřeném systému pouze jediný druh (viz (8) a (9)), zatímco ostatní postupně vyhynou. Lze ukázat (viz [2]), že půjde o druh s nejvyšší hodnotou e_i , tj. ten, který je schopen nejefektivněji zpracovávat stravu ve prospěch přírůstu vlastních jedinců.

Zkomplikujme nyní situaci možností, že se v určitém čase v systému objeví mutant, tj. jedinec, jehož e_i se bude lišit od e_i jeho rodiče a půjde tak vlastně podle našich kritérií o zakladatele nového druhu. V modelu předpokládáme dobu mezi výskytem dvou mutantů náhodnou s rovnoměrným rozdělením pravděpodobnosti, což sice neodpovídá reálnému prostředí, nicméně zřejmě samotná perioda výskytu mutantů nemá na stav systému z dlouhodobého hlediska vliv.

Hodnotu e_i mutovaného potomka pak budeme uvažovat v intervalu $\left(\frac{e_j}{2}; \frac{3 \cdot e_j}{2}\right)$, kde e_j odpovídá hodnotě efektivity přeměny substrátu na živou hmotu jeho rodiče. Rozdělení pravděpodobnosti využijeme pro jednoduchost opět rovnoměrné. V modelu budeme uvažovat stejnou pravděpodobnost zplození mutantu pro všechny žijící jedince, tj. pravděpodobnost vzniku mutantu i -té populace bude přímo úměrná velikosti této populace.

2. Implementace modelu v prostředí Maple

Základem modelu v softwarovém prostředí Maple 16 je cyklus, jehož jedna iterace odpovídá období mezi výskytem dvou mutantů. Časová délka období reprezentovaného iterací je proměnlivá. Jednotlivé populace jsou v modelu reprezentovány řádky populační matice *Matice*, jejíž první sloupec udává pro každou populaci efektivitu e_i , druhý sloupec pak odpovídá počtu jedinců dané populace.

celkemcyklu := 250;

a := 2.5;

b := 0.8;

d := 12;

m := 0.1;

z := 0.1;

cas := 0

q := proc (x) options operator, arrow; a*x/(1+a*b*x) end proc

Před spuštěním cyklu je vygenerována dvouřádková matice se dvěma náhodnými populacemi, které vstupují do prvního období. Délka období je volena náhodně z uživatelem zadaného intervalu. Na začátku každého období je provedena kontrola všech řádků matice a „neúspěšné“ populace s počtem jedinců nižším než 1 (tj. populace nesplňující zadanou podmínku pro další rozmnožování)⁵ jsou z modelu odstraněny. Následně probíhá vygenerování náhodného mutanta, který je zapsán jako nový řádek do matice (tj. populace s jediným jedincem).

```
# Matice se dvema nahodnymi populacemi;
h1 := rand(1 .. 10);
h2 := (1/100)*rand(1 .. 100);
Matice := Matrix(2, 2, [[h2(), h1()], [h2(), h1()]]);

# Hlavni cyklus;
for cyklus from 1 to celkemcyklu do

  # Cyklus pro odstraneni populaci s mene nez 1 jedincem;
  v:=0:
  for i from 1 by 1 to RowDimension(Matice) do
    if (Matice[i-v,2]<1) then
      Matice:=DeleteRow(Matice,i-v);
      v:=v+1
    end if:
  end do:

  # Cyklus pro vznik jednoho jedince s nahodnou mutaci;
  jedincu:=add(Matice[k,2],k=1..RowDimension(Matice));
  h:=(rand(1..10000* floor(jedincu)))/(10000):
  budemutant:=h():
  soucet:=0:
  j:=0:
  for j from 1 by 1 while soucet<budemutant do
    soucet:=soucet+Matice[j,2]:
  end do:
  mutant:=Matice[j-1,1]:

  # Mutace a zapis mutanta do matice:
  h:=(rand(0..100)-50)/(100):
  mutant:=mutant+mutant*h():
  Matice(RowDimension(Matice)+1,1..):=<mutant,1>:
```

V dalším kroku je pomocí vnořeného cyklu vytvořena soustava $n + 1$ diferenciálních rovnic postupně pro jednotlivé populace z matice a navíc pro celkové dostupné množství substrátu v systému, z velikosti populací a proměnné Y , udávající množství dostupného substrátu po skončení předešlého období jsou vygenerovány počáteční podmínky pro řešení a následně je soustava řešena pomocí numerické varianty příkazu *dsolve*. Řešení je z důvodu omezení časové náročnosti hledáno pouze pro období mezi výskyty mutantů, po jehož uplynutí se změní počáteční podmínky.

```
# Vytvoreni soustavy diferencialnich rovnic;
```

⁵ Pro snadnější průběh výpočtů budeme počet jedinců populace považovat za spojitou veličinu a budeme tedy pracovat i s neceločíselnými hodnotami. Pro názornější představu lze hodnotu P_i považovat např. za celkovou hmotnost všech jedinců populace, což odpovídá rovněž procesu přeměny potravy na živou hmotu (tj. rodičů i potomků současně).

```

rovnice[0]:=diff(S(t),t)=d-z*S(t)-
q(S(t))*add(P[i](t),i=1..RowDimension(Matice)):
for i from 1 by 1 to RowDimension(Matice) do
    rovnice[i]:=diff(P[i](t),t)=-m*P[i](t)+q(S(t))*P[i](t)*Matice[i,1]:
end do:
soustava:=seq(rovnice[i],i=0..RowDimension(Matice)):

# Vytvoreni pocatecnich podminek;
if cas=0 then
    podminka[0]:=S(0)=1: else podminka[0]:=S(cas)=Y:
end if:
for i from 1 by 1 to RowDimension(Matice) do
    podminka[i]:=P[i](cas)=Matice[i,2]:
end do:
podminky:=seq(podminka[i],i=0..RowDimension(Matice)):

# Cas do pristi mutace;
minule:=cas:
h:=rand(1..40):
cas:=cas+h(1..40):

# Reseni a vykresleni do grafu;
reseni:=dsolve({soustava,podminky},numeric, maxfun=10000000):
krivky:=seq([t,P[i](t)],i=1..RowDimension(Matice)):
graf[cyklus]:=odeplot(reseni,([krivky],minule..cas)):
assign(reseni(cas)):

# Zapis novych hodnot do matice;
for i from 1 by 1 to RowDimension(Matice) do
    Matice[i,2]:=P[i](t):
    unassign('P[i](t)'):
end do: Y:=S(t):
unassign('t'):
unassign('P'):
unassign('S'):

end do:

```

Vzhledem k tomu, že S konverguje v období mezi dvěma mutacemi (jak bylo výše ukázáno) ke konstantě, konverguje podle definice ke konstantní hodnotě rovněž funkce příjmu potravy q a tedy můžeme provést následující výpočet:

$$\begin{aligned} \frac{dP_i}{dt} &\xrightarrow{t \rightarrow \infty} P_i \cdot e_i \cdot q - mP_i \\ P_i &\xrightarrow{t \rightarrow \infty} e^{P_i \cdot (qe_i - m)} \end{aligned} \quad (8)$$

Tedy zřejmě velikost populace v rámci období mezi mutacemi konverguje k monotónnímu chování. Předpokládejme nyní:

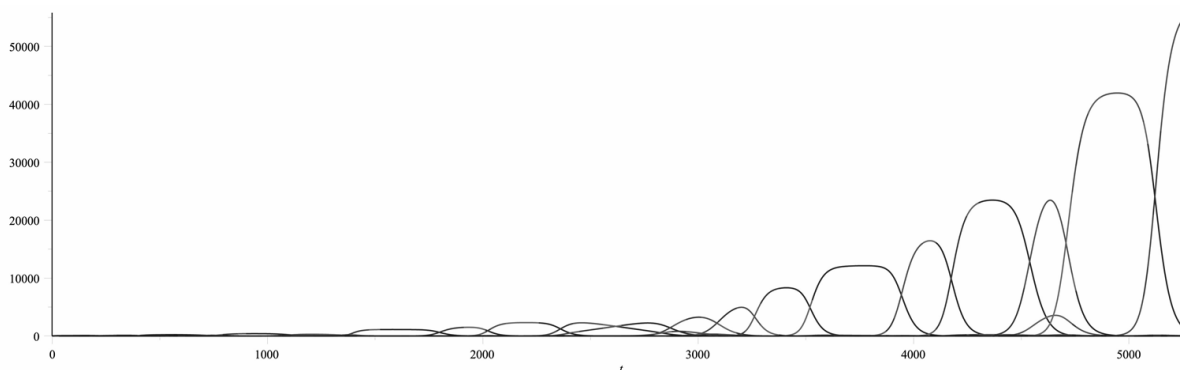
$$\begin{aligned} P_i &\xrightarrow{t \rightarrow \infty} konst. \Rightarrow P_i \cdot e_i \cdot q - mP_i \xrightarrow{t \rightarrow \infty} 0 \\ P_i &\xrightarrow{t \rightarrow \infty} 0 \vee e_i \xrightarrow{t \rightarrow \infty} \frac{m}{q_i} \end{aligned} \quad (9)$$

Protože pro libovolné populace $i, j: i \neq j$ platí $e_i \neq e_j$, konverguje v období mezi dvěma mutacemi velikost všech populací, případně vyjma populace s nejvyšším e_i , k nule.

Příkazem `display(seq(graf[i], i = 1 .. celkemcyklu))` si nyní můžeme zobrazit přes sebe graf vývoje populací ve všech obdobích (obr. 1) a na první pohled je zřejmé, že skutečně v každém období mezi dvěma mutacemi je jediná populace, jejíž velikost roste „na úkor“ ostatních populací. Ve variantě bez mutací by se po určité době velikost této populace ustálila na hodnotě úživnosti prostředí, zatímco velikost konkurenčních populací by klesla k nule. Náhodně se vyskytující mutace však do systému vnáší dynamické chování, neboť v případě, kdy se objeví mutant s e_i vyšším, než bylo dosavadní nejvyšší e_j , rozšíří se tato i -tá populace a postupně zapříčiní vyhynutí všech populací včetně j -té, pokud není ještě dříve vystřídána jinou populací, vzešlou z mutantu s vyšším e_k .

Vzhledem k tomu, že v této variantě modelu není výše efektivity přeměny potravy na živou hmotu e ničím omezena, dochází k tomu, že se neustále objevují mutanti s vyšší hodnotou e_i , která může v rozporu s realitou překročit dokonce hodnotu 1 a růst neomezeně do nekonečna. Velikost momentální největší populace se tak bude v modelu pohybovat v čase po exponenciální křivce donekonečna. Takový proces je samozřejmě zcela nerealistický a proto je třeba model upravit.

Zvyšování efektivity přeměny potravy na živou hmotu vlastního druhu se v modelu ukazuje jako správná strategie vedoucí k přežití. Je ovšem patrné, že s rostoucím e_i vítězná populace dochází k poklesu substrátu v prostředí, neboť vítězná populace ostatní jednoduše „vyhladoví“. V našem nerealistickém případě z dlouhodobého hlediska dochází až ke konvergenci S k nule. Tento pokles úživnosti a obecně evoluční zhoršení podmínek prostředí na ty nejhorší možné se nazývá pesimizací prostředí.



Obr. 1: Exponenciální růst populací v modelu ve variantě bez omezení efektivity e_i přeměny substrátu na živou hmotu.

3. Varianta s omezením efektivity

Fyzikální podstata živých tvorů neumožňuje růst efektivity e_i nad hodnotu 1 (ve skutečnosti je číslo vzhledem k chemickým procesům v organismu samozřejmě ještě nižší), vlastní fyziologie organismů však znamená ještě výraznější pokles efektivity, neboť v reálném prostředí nemůže věnovat žádný jedinec maximum prostředků na hospodárné nakládání se získanou energií (tj. vlastní růst a reprodukci), ale musí rovněž zajistit řadu dalších procesů, mezi které kromě jiných strategických činností patří samotný příjem potravy.

V předchozí zjednodušené variantě modelu jsme předpokládali závislost množství přijaté potravy q pouze na množství substrátu S v prostředí, zbývající parametry a a B_i jsme považovali za konstantní pro všechny populace. Definujme nyní vztah mezi efekivitou přeměny substrátu na živou hmotu e_i a množstvím potravy, kterou je jedinec schopen přijmout za jednotku času q_i , který bude reprezentovat strategii druhu ve smyslu rozhodování mezi investicí energie do většího množství získané potravy nebo lepším využitím potravy pro vlastní růst a reprodukci.

Nahradíme tedy v rovnici logistické křivky (2) konstantu B_i kvadratickou funkcí e_i takto:

$$B_i = e_i^2 + c \quad (10)$$

tj. s rostoucí efektivitou $e_i > 0$ bude injektivně klesat množství substrátu q_i , které je jedinec daného druhu schopen zkonzumovat. Tato vazba zamezí nekontrolovanému růstu e_i a přiblíží model realitě. Optimální strategií pak bude dosáhnout takového e_i , aby byla hodnota $e_i \cdot p_i$ (viz (3)) maximální, tj:

$$q_i \cdot e_i = \frac{e_i \cdot a \cdot S}{1 + (e_i^2 + c) \cdot a \cdot S} \quad (11)$$

a po derivaci podle e_i :

$$\frac{a \cdot S \cdot (1 - a \cdot S \cdot e_i^2 + a \cdot S \cdot c)}{(1 + a \cdot S \cdot e_i^2 + a \cdot S \cdot c)^2} = 0 \quad (12)$$

Po úpravě, protože jmenovatel je nenulový, získáváme vztah pro optimální $\hat{e}$:

$$\hat{e} = \sqrt{\frac{1}{a \cdot S} + c}. \quad (13)$$

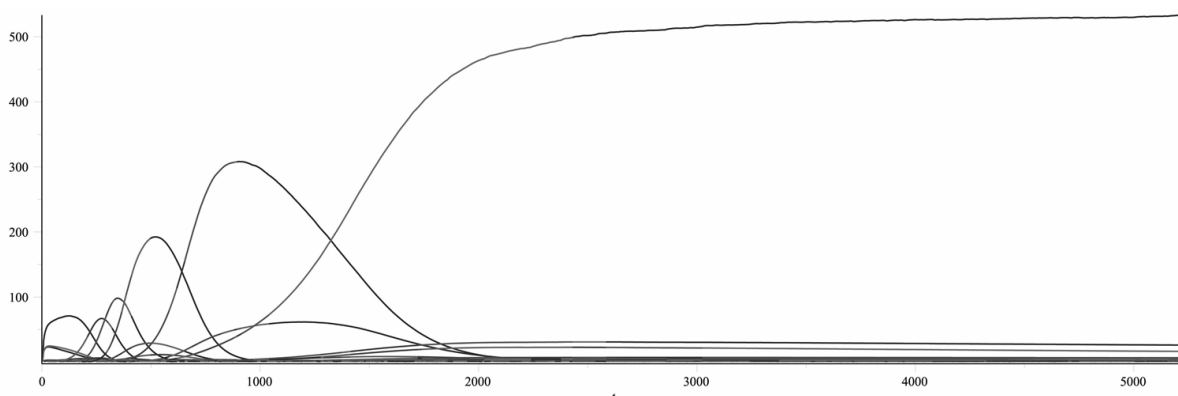
Ve zdrojovém kódu Maple stačí změnit definici funkce q_i na funkci dvou proměnných S a e_i a upravit cyklus určující tvar diferenciálních rovnic:

```
q := proc (x, y) options operator, arrow; a*x/(1+a*(y^2+c)*x) end proc
# Vytvoreni soustavy diferencialnich rovnic;
rovnice[0] := diff(S(t), t) = d-z*S(t)-add(q(S(t), Matice[i, 1])*P[i](t), i
= 1 .. RowDimension(Matice));
for i to RowDimension(Matice) do
    rovnice[i] := diff(P[i](t), t) = -m*P[i](t)+q(S(t), Matice[i,
1])*P[i](t)*Matice[i, 1]:
end do;

soustava := seq(rovnice[i], i = 0 .. RowDimension(Matice))
```

Podle očekávání je po vykreslení grafu jasně patrný nárůst celkového e_i momentálně nejúspěšnějších populací, který však zpomaluje až ke konstantní hodnotě odpovídající úživnosti prostředí pro populaci s hodnotou blízkou $\hat{e}$ podle rovnice (13). Z dlouhodobého hlediska by za neměnných podmínek v systému opravu po řadě mutací převládla jediná populace s $e_i = \hat{e}$.

Protože $\hat{e}$ je nejvyšší ze všech dlouhodobě možných efektivit přeměny substrátu na živou hmotu, dojde postupem času v prostředí opět k poklesu dostupného množství potravy na nejnižší možnou mez – vzájemný souboj populací o zdroje tedy stejně jako v předchozím případě nastoluje ty nejhorší podmínky pro přežití, tj. dochází k evoluční pesimizaci.



Obr. 2: Omezený růst populací v modelu s omezením efektivity e_i přeměny substrátu na živou hmotu. Po dosažení hodnoty $\hat{e}$ se růst zastavuje na hodnotě úživnosti prostředí.

4. Závěr

Ve variantě modelu, kdy není efektivita přeměny substrátu na živou hmotu e shora omezena, dochází vlivem prostředí za vzniku náhodných mutací k postupnému zvyšování e přibližně exponenciálním tempem, přičemž každá populace po určité době vyhyne, protože úživnost prostředí bude snížena populací s vyšší efektivitou e pod přípustnou mez.

Pokud je do modelu vložena zpětná vazba omezující růst efektivity e jeho klesající úměrností vůči množství zkonsumovaného substrátu q , dochází postupnými mutacemi k přibližování maximální efektivity e k mezní hodnotě úživnosti prostředí $\hat{e}$, která je daná fyziologickými vlastnostmi organismů. V tomto případě platí odvozené vztahy dokazující, že nejvýše jedna populace v dlouhodobém měřítku nevyhyne a obsadí celý systém. Pro tuto populaci platí podmínka $e = \hat{e}$, tj. není ohrožena žádným mutantem.

Vlivem evolučních změn vedoucích k postupnému hynutí populací s nižší hodnotou e než mají jejich konkurenční mutanti, dochází k tzv. pesimizaci prostředí, tj. snížení úživnosti na nejmenší přístupnou mez. V krajním případě může tento jev vést až k úplnému vyhynutí všech populací.

5. Literatura

- [1] Hřebíček, J., Pospíšil, Z., Urbánek J.: Úvod do matematického modelování s využitím Maple. Akademické nakladatelství CERM, s.r.o., Brno, 2010.
- [2] Diekmann, O. A beginner's guide to adaptive dynamics. Mathematical modelling of population dynamics, Banach Center Publications, vol. 63. Institute of Mathematics Polish Academy of Sciences, Warszawa, 2004.

System pro optimalizaci medicínského kurikula

Martin Komenda

Institut biostatistiky a analýz, Kamenice 126/3, 625 00, Brno
Fakulta informatiky, Botanická 68a, 602 00, Brno
komenda@iba.muni.cz

Abstrakt

V příspěvku je představena zcela nová a původní metodika optimalizace medicínských osnov v rámci terciárního vzdělávání s využitím outcome-based přístupu a aplikace moderních informačních a komunikačních technologií. Existující publikovaná řešení se zaměřují na kurikulum pouze z určitého pohledu a nabízí agendu spolu s vybranými funkcemi, které se snaží zpřístupnit v přehledné formě studentů a pedagogů dané instituce. Nicméně komplexní nástroj, který by současně zahrnoval všechny prvky spojené s globální optimalizací kurikula včetně detailního parametrického popisu až na úroveň samotných tematicky ucelených bloků výuky prozatím neexistuje. Záměrem je tedy navrhnout zcela nový webové orientovaný nástroj včetně propracované metodiky, který podpoří optimalizaci medicínského kurikula s využitím outcome-based přístupu.

Abstract

In the contribution a brand new and original medical curriculum optimization methodology within tertiary education will be described by adopting an outcome-based approach and applying modern IT and communication technologies. Existing solutions that have been published are focused on the curriculum only from a certain perspective, offering the agenda together with selected functionalities and making an effort to provide them to students and teachers of the respective institution in a transparent format. However, there still does not exist a complex instrument that would cover all elements connected with global curriculum optimization, including a detailed parametric description down to the level of learning units. The aim is to create a new web-oriented tool included advanced methodology, which provides optimization of medical curriculum using learning outcome approach.

Klíčová slova

vzdělávání, optimalizace, medicínské kurikulum, webové technologie

Key words

education, optimization, medical curriculum, web-based technology

1. Úvod

Studium mladých lékařů na vysoké škole je oproti jiným oblastem do jisté míry specifické. Důvodem je především fakt, že jejich další uplatnění v praxi nepřipouští nedostatky ve znalostech nabitých po dobu studií a veškeré chyby mohou mít mnohdy fatální následky. Potřeba garantovaného a kvalitního vzdělání, které zahrnuje předepsané osnovy pokrývající odpovídající rozsah výstupních znalostí a dovedností současně požadovaných v navazující praxi, je stále hlasitěji akcentována. Lékařské fakulty sestavují svá kurikula tak, aby komplexně pokryly nezbytné nároky nutné pro další uplatnění svých absolventů. Ti musí splnit všechny předepsané povinnosti v podobě úspěšného ukončení povinných a povinně volitelných kurzů včetně státní závěrečné zkoušky a poté se připravují na získání odborné způsobilosti k výkonu povolání lékaře v podobě atestace.

Správně sestavené a vyvážené kurikulum napříč medicínskými obory je nezbytným předpokladem pro výchovu mladých lékařů. Vhodná kombinace teoreticky zaměřených kurzů s klinickou výukovou základnou je bez pochyb klíčem k optimálnímu návrhu studijních osnov. Lékařské fakulty nejčastěji poskytují vzdělání v 6letých magisterských studijních programech Všeobecné lékařství, které jsou zakončeny udělením titulu MUDr. (doktor všeobecného lékařství). Podobně jako v jiných oborech a na jiných fakultách se i v medicíně ukazuje, že přehled o struktuře a obsahu výuky není zcela ideální a mnohdy nastává situace, že překryv teoretických a klinických předmětů mezi sebou i napříč je příliš velký nebo naopak nedostačující. Při rychlém rozvoji moderních informačních technologií a oblibě

pohodlné práce na Internetu se nabízí prostor pro návrh technologie, která by nepřehlednost výukových osnov nejen eliminovala a významně by tak napomohla ke zkvalitnění celé výuky.

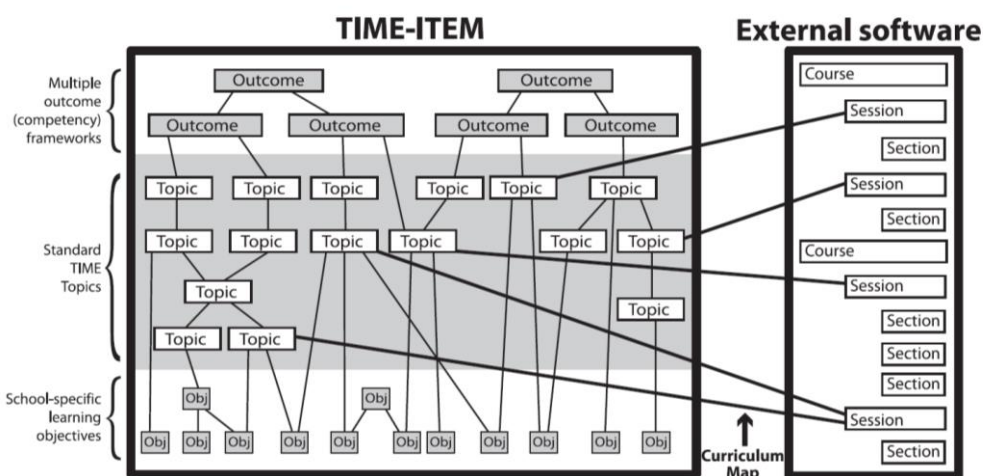
Seznam jednotlivých povinných, povinně volitelných a volitelných předmětů, jejich anotace a osnovy jsou studentům a pedagogům běžně dostupné, typicky v lokálních systémech pro správu výuky. Nicméně různá úroveň detailu i styl popisu bez jakékoli míry standardizace nebo parametrizace vede k tomu, že se celková přehlednost a srozumitelnost obzvlášť při snaze dohledat informace o celém studiu vytrácí. Je tak velmi obtížné podívat se na celý obor, specializaci nebo studium ze širší perspektivy a mít možnost napříč kurikulem snadno vyhledávat a tedy se i celkově orientovat v tom, co se kde a jak vlastně učí. Pro studenta by takový pohled znamenal jasnou informaci o tom, jaké znalosti je nutné si za šestiletou dobu studia osvojit, s jakými tématy se ve výuce seznámí, které oblasti budou akcentovány opakovaně a které konkrétní předměty jsou s danou problematikou spojeny. Pro pedagogy a vedení fakulty by výše zmíněný přehled reprezentoval praktický pohled na výuku, srozumitelně by demonstroval kdo, co a v jakém rozsahu učí, zda se vyučující tematicky překrývají, zda je tento překryv žádoucí či nežádoucí, co se vyučuje v klinických a teoretických oborech, zda je celkové rozvržení výuky správné nebo naopak, jestli není zapotřebí restrukturalizace.

2. Přístup založený na výstupech z učení

Kontinuální představování nových výukových technik a exponenciální růst uživatelů používající každodenně Internet předesílá značný potenciál ke změně terciálního vzdělávání [1]. Na bergenské vládní konferenci v květnu 2005 byla diskutována témata spojená s reformami v oblasti školství. Evropské systémy tak prochází radikální restrukturalizací, která vychází ze záměrů definovaných v Bergenu. Jako základní stavební kámen pro inovaci studijních programů byl stanoven přístup založený na tzv. výstupech z učení (outcome-based approach). Použití výstupů z učení významně ovlivňuje fundamentální paradigma pro návrh kurikula mnoha evropských institucí nabízející vysokoškolské vzdělání. S. Adams definoval výstupy z učení jako soubor znalostí, dovedností, schopností, přístupů a porozumění, které by si student měl osvojit v rámci daného výukového celku [2]. Tento koncept je v současné době hojně využíván s jednoznačným cílem zpřehlednit vlastní výuky a jasně definovat soubor výstupních požadavků kladených na studenty.

2.1 Systémy pro správu kurikula

Problematikou inovace kurikula se nejen v oblasti lékařských a zdravotnických oborů zabývají různé akademické instituce. Vzhledem k požadavku podpory přístupu založeném na výstupech z učení, vznikly



Obr. 1 – Struktura aplikace TIME. Zdroj [3]

a byly publikovány nové systémy umožňující efektivnější správu kurikula. Níže jsou krátce představeny dva nejvýznamnější. Popis, vzájemné provázání a grafické znázornění závislostí mezi jednotlivými moduly, tématy a cíli výuky, poskytuje uživatelům webová aplikace TIME [3]. Jako nedostatek lze v tomto případě označit fakt, že aplikace je určena pro generování indexů popisující medicínské vzdělávání a pro samotnou vizualizaci veškerých nadefinovaných vztahů je nutný export do externího prostředí, které disponuje požadovanou funkcionalitou. Podobně interaktivní prostředí ActiveCC Web [4] umožňuje sofistikovaný popis kurikula a nabízí provázání konkrétních modulů. Navíc určuje, které moduly musí student úspěšně absolvovat před vstupem do modulů navazujících. Přehledná diagramatická struktura je dostupná pouze na úrovni modulu, ale nelze se zanořit do podrobnější úrovně, která je však pro cílovou skupinu velmi užitečná. Detailní náhled umožní pedagogům snadno najít konkrétní obsahové překryvy či naopak nedostatky v osnovách výuky a celkově se tak ve výuce daného oboru lépe zorientovat. ActiveCC web přináší přidanou hodnotou pro studenty ve formě srozumitelného přehledu o požadavcích na absolventa (výstupech z učení) a provázání na reálné kurzy, kde se daná problematika probírá, případně na další související oblasti výuky.

Existující publikovaná řešení se zaměřují na kurikulum pouze z určitého pohledu a nabízí agendu spolu s vybranými funkcemi, které se snaží zpřístupnit v přehledné formě studentům a pedagogům dané instituce. Nicméně komplexní nástroj, který by současně zahrnoval všechny prvky spojené s globální optimalizací kurikula včetně detailního parametrického popisu až na úroveň samotných tematicky ucelených bloků výuky, tzv. výukových modulů včetně vazby na výukové objekty prozatím neexistuje.

3. Projekty na podporu medicínského vzdělávání

3.1 Vzdělávací síť MEFANET

Jedinečný mezifakultní projekt MEFANET (MEDical FACulties NETwork) se zaměřuje na podporu rozvoje výuky lékařských a zdravotnických oborů s využitím moderních informačních technologií. Od roku 2007, kdy tato vzdělávací síť vznikla, neustále posiluje svou pozici a počet zainteresovaných členských akademických institucí narůstá [5]. Primární cíl se od založení nezměnil a spočívá ve vývoji sofistikovaných řešení, která umožňují široké akademické obci snadno a rychle přistupovat k elektronickému výukovému obsahu napříč celou sítí lékařských fakult. To vše se děje s plným uvědoměním a respektem vůči nezávislosti a samostatnosti zapojených fakult. V současnosti jsou do projektu zapojeny nejen všechny české a slovenské lékařské fakulty, ale také další fakulty a akademické subjekty, jejichž zaměření se více či méně dotýká vzdělávání v medicínské oblasti.

3.2 Projekt OPTIMED

Primární snahou projektu OPTIMED (Optimization of Medical Education) je komplexní inovace systému výuky všeobecného lékařství na LF MU a posílení výuky orientované na řešení problému v souladu s uplatněním absolventa v oblasti klinické i akademické. Stěžejními prvky projektu jsou:

- Horizontální inovace všech vyučovaných předmětů s využitím outcome-based přístupu a nástrojů dostupných v rámci vyvíjené platformy (prohlížeč výstupů z učení, registr výukových jednotek, repozitář výukových objektů, reportovací nástroje).
- Vertikální propojení výuky na ose: vstupní znalosti studentů medicíny - teoretické a preklinické znalosti - klinické znalosti a dovednosti – schopnosti lékaře – absolventa po nástupu do praxe.

OPTIMED se tak opírá o vytvoření inovativní dynamické platformy, která bude usnadňovat studentům i vyučujícím orientaci ve výuce a ve svém důsledku zefektivňovat znalosti a dovednosti studentů pro praxi. Klíčovým parametrem systému je jeho dynamičnost, tedy schopnost absorbovat nové poznatky v medicíně a racionálním způsobem je propojit s výukou orientovanou na pacienta [6]. Primární úkolem není radikální změna výuky, ale s pomocí vhodně zvolených ICT důkladně zmapovat současný stav na LF MU a umožnit přehlednou orientaci napříč výukou.

4. Systém pro optimalizaci kurikula

Inovativní systém je založen na zcela nové a původní metodice optimalizace medicínských osnov v rámci terciárního vzdělávání s využitím outcome-based přístupu a aplikace moderních informačních a komunikačních technologií. Hlavní přínosy a benefity vyvíjeného řešení spočívají v návrhu konceptu postaveného na níže uvedených funkcionalitách, které však doposud v podobných řešeních nebyly implementovány:

- i. Návrh původní metodiky optimalizace medicínského kurikula napříč klinickými a teoretickými obory, která je založena na propracované parametrizaci studijních osnov s využitím outcome-based paradigmatu.
- ii. Návrh databázového uspořádání metadat popisujících kurikulum všeobecného lékařství nezávisle na následné implementaci.
- iii. Vývoj webové orientované platformy vycházející s výše uvedeného databázového modelu, která umožní reálné využití popsané metodiky při optimalizaci vybraného studijního oboru. Tento modulárně strukturovaný systém poskytne akademické veřejnosti intuitivní kolaborativní prostředí pro snadnou tvorbu a prohlížení obsahu kurikula včetně efektivního prohledávání a reportovacích nástrojů pro interaktivní grafické zobrazení vazeb, vztahů a podobností mezi dostupnými daty.
- iv. Integrace tezauru MeSH [7] pro standardizovanou práci s klíčovými slovy. Jelikož se jedná o optimalizaci konkrétního oboru vyučovaného v českém jazyce, byl zvolen slovník, který je standardizován v anglickém jazyce a současně disponuje českým každoročně aktualizovaným překladem.
- v. Návrh propracovaného systému přístupových oprávnění v souladu s představenou metodikou. Systém rozlišuje uživatelské role podle míry zapojení do optimalizačního procesu. Současně bude integrovaný centrální ověřovací mechanismus prostřednictvím technologie Shibboleth (8), který umožňuje snadné rozšíření přístupu na úroveň uživatele vzdělávací sítě MEFANET (9), tedy studenta nebo pedagoga jedné z jedenácti českých a slovenských lékařských fakult.
- vi. Integrace vybraných funkcí morfologického analyzátoru českého jazyka Majka [10] pro zefektivnění vyhledávání a analytického zpracování.
- vii. Systémové zpracování metadat popisujících doporučené tištěné i elektronické studijní podklady s přímou vazbou na konkrétní výstupy z učení a výukové jednotky. Vzhledem k dalšímu rozvoji e-learningové agendy na LF MU je pedagogům umožněno vznést nové požadavky a definovat tak výukové materiály, které v současnosti ve své výuce postrádají.
- viii. Vytvoření unikátního slovníku významných pojmů pro studium všeobecného lékařství, který se bude generovat z definovaných atributů datové věty jednotlivých výstupů z učení.
- ix. Vazba na lokální VLE (Virtual Learning Environment) propojením s edukačním portálem vzdělávací sítě lékařských fakult MEFANET a Informačním systémem pro správu výuky na Masarykově univerzitě [11].

Výše uvedené cíle se přímo dotýkají projektu OPTIMED. Metodika včetně vyvinuté technologické platformy bude v praxi použita zkušenými pedagogy a odbornými garanty výuky v rámci obsahového auditu výuky oboru Všeobecné lékařství. V současnosti je vedením Lékařské fakulty delegováno více než 250 vyučujících, kteří zasáhnou v různých rolích do procesu optimalizace a zpřehlednění vysokoškolského medicínského studia, a kteří poskytnou zpětnou vazbu k navržené inovativní technologii. Konkrétní zaměření na oblast medicíny nemá z globálního pohledu vliv na obecné použití. S mírnými úpravami bude možná aplikace navrženého modelu optimalizace na jakýkoli obor studia na vysoké škole. Předkládané řešení obsahuje určitá specifika v podobě integrace biomedicínského slovníku a přímých vazeb na existující virtuální systémy pro podporu studia zaměřené cíleně na medicínské vzdělávání.

4.1 Architektura systému

Architektura platformy se skládá ze dvou částí, kterými jsou FrontOffice a BackOffice. FrontOffice reprezentuje rozhraní použité pro prezentaci obsahu koncovému uživateli. Obsah této sekce bude volně dostupný pouze pro cílovou skupinu uživatelů, tedy pro studenty a pedagogy LF MU s možným rozšířením po ukončení vertikální optimalizace pro členy akademické obce napříč vzdělávací

síť MEFANET. I z tohoto důvodu bude využito jednotného rámce pro ověřování identity uživatele, který poskytuje česká akademická federace identit eduID.cz [12]. Konkrétně se jedná o zprovoznění technologie Shibboleth, která zajistí uživatelskou autentizaci prostřednictvím domovské instituce. V okamžiku přístupu ke chráněnému obsahu je uživatel automaticky přesměrován na svého poskytovatele identity. Zde prokáže svou totožnost prostřednictvím lokálních přihlašovacích údajů a následně je vrácen zpět na původní server. BackOffice je druhou částí platformy, která slouží administrátorům a redaktorům, jakožto rozhraní pro vkládání a editaci statického obsahu.

Protože platforma bude sloužit uživatelům jako nástroj pro přehled výukových osnov, bude pozornost zaměřena také na efektivní a rychlé vyhledávání. Ve spolupráci s výzkumným záměrem Fakulty informatiky Masarykovy univerzity se dále rozvíjí funkcionality morfologického analyzátoru pro český jazyk Majka [10], [13]. Cílem je integrace skriptu, který bude zpracovávat jednotlivé vstupní řetězce do kořenových tvarů a poté ukládat tuto informaci spolu s vazbou na související metadata do databáze. V případě, že uživatel zadá hledaný výraz, s využitím nové funkcionality budou jako výsledky zobrazeny všechny slova spojená s kořenem hledaného výrazu. Vzhledem k tomu, že současné fulltextové vyhledávání neposkytuje pro český jazyk spolehlivý způsob, jak pracovat se slovními tvary, je zvolený způsob zpracování zajímavým a přínosným řešením. Také při implementaci analytických a vizualizačních nástrojů budou hrát kořenové tvary slov významnou roli.

Dalším prvkem je použití české mutace tezauru Medical Subject Headings (MeSH) s cílem standardizovat klíčová pojmy spojené s edukačním obsahem výukových jednotek, jejichž popis se ukládá do databáze. Samotná klíčová slova jsou definována a strukturována v mnoha podobách a potřeba sjednocení s ohledem na mezinárodní rámec stále vzrůstá. Biomedicínský slovník MeSH je v anglickém jazyce vydáván od roku 1960 americkou Národní lékařskou knihovnou. Český překlad tohoto tezauru vytváří Národní lékařská knihovna ČR, která vydává pravidelné roční aktualizace. Slovník obsahuje 26 142 hesel s více než 54 000 odkazy (7). Po smluvně podložené dohodě s NLK ČR bude MeSH využit také pro účely nově vznikající platformy pro optimalizaci kurikula. Hlavním požadavkem při integraci standardizovaného slovníku byla pravidelně vydávaná doplnění české mutace, což MeSH jako jediný z dostupných řešení splňuje. Nyní se neuvažuje o rozšíření do jiných jazykových variant, ale obecně by případná změna neměla vzhledem k navržené struktuře znamenat přílišné komplikace.

Kromě klíčových slov definovaných mezinárodně uznávaným a standardizovaným formátem, se v popisu výukových jednotek vyskytují tzv. významné pojmy, pro které žádný slovník neexistuje. Databázi důležitých pojmů obsažených ve výuce oboru Všeobecné lékařství, se kterými se studenti v průběhu šesti let setkají a které by měli absolventi znát, budou plnit samotní pedagogové a garanti studia. Vznikne tak zcela unikátní slovník významných hesel medicínského studia, který poskytne další z řady způsobů, jak se orientovat ve složitém a obsáhlém obsahu kurikula. Pro alespoň částečnou standardizaci (kvazistandardizaci) bude uživatelům nabízen při vkládání slov našeptávač, který zobrazí termíny korespondující s již napsanou částí textu, konkrétně existující slova či odpovídající položky, které již definovali předešlí autoři. Po finálním pročištění, které však nelze plně zautomatizovat, lze jako jeden z výstupů a tedy i benefitů označit doposud neexistující slovník významných pojmů z oblasti výuky oboru Všeobecného lékařství.

Platforma bude mimo data popisující kurikulum obsahovat také přímé vazby na virtuální systémy podpory výuky, které nabízí uživateli jak další potřebné informace o studiu, tak i samotné elektronické studijní materiály. Jedná se zejména o prozatím jednosměrné propojení s Informačním systémem Masarykovy univerzity a provázání s edukačním portálem vzdělávací sítě českých a slovenských lékařských fakult MEFANET. Vybudování a následné propojení fundamentu v podobě nového repozitáře znovupoužitelných výukových objektů, v odborné literatuře známých jako RLO (reusable learning objects), závisí na počtu a zvoleném formátu nově vytvořených elektronických studijních opor. Navržené řešení bude na integraci RLO uložistiště připraveno včetně aplikace vybraného metadatového standardu.

Jednou z nejprínosnějších částí projektu je vizualizace dat, která byla parametricky zadána odbornými guaranty výuky. Nezávislé moduly budou umožňovat různé pohledy na kurikulum. Moduly Prohlížeč výstupů z učení a Registr výukových jednotek budou založeny na komponentě datového gridu, který

optimálně zpřístupní uživateli data v přehledné formě včetně možnosti aplikovat pokročilé vyhledávání a filtrování dle zvolených atributů. Modul reportovací nástroje bude zahrnovat grafické znázornění dostupných dat. Aplikace metod data miningu v kombinaci s analytickým zpracováním poskytnou podklad pro vizualizaci informačně přínosných a hodnotných vazeb napříč kurikulem, které nabídnou uživateli originální pohled na výuku v daném oboru ve formátu interaktivních sémantických sítí a grafů. Nedílnou součástí bude zkoumání podobnosti jednotlivých kurzů a na nižší úrovni i výukových jednotek podle terminologického obsahu.

5. Závěr

Pro realizaci navrženého řešení je nutné nejprve parametricky popsat edukační proces a stanovit, které informace a v jakém objemu jsou pro následnou revizi a další zpracování důležité. Navržený konceptuální model tvoří základ pro možnou implementaci v praxi. Podrobně popisuje všechny zainteresované subjekty a definuje vztahy mezi nimi tak, aby bylo možné všechny zvolené vazby později přehledně a jasné zobrazit. Celý proces optimalizace je rozdělen do několika na sobě závislých fází umožňující efektivně zmapovat kurikulum daného oboru nebo specializace studia. Pro každou fázi je charakteristický doporučený nástroj, který cílovému uživateli – pedagogovi – poskytne možnost intuitivní práce při vytváření/prohlížení určeného obsahu. Veškeré nástroje jsou vyvíjeny na základě reálných požadavků a jsou praktickou ukázkou implementace metodiky s nasazením vhodných informačních technologií. Na rozdíl samotné platformy je aplikace metodiky zcela nezávislá na zvolených ICT stejně jako na oblasti zaměření (volba studijního oboru, který projde procesem optimalizace).

6. Použitá literatura

- [1] Harden RM, Hart IR. An international virtual medical school (IVIMEDS): the future for medical education? *Medical Teacher*. 2002 May;24(3):261–7.
- [2] Adam S. Using learning outcomes - A consideration of the nature, role, application and implications for European education of employing “learning outcomes” at the local, national and international levels. University of Westminster, Edinburgh; 2004.
- [3] Willett TG, Marshall KC, Broudo M, Clarke M. TIME as a generic index for outcome-based medical education. *Medical Teacher*. 2007 Jan;29(7):655–9.
- [4] Kabicher S, Derntl M, Motschnig-Pitrik R. ActiveCC - A Collaborative Framework for Supporting the Implementation of Active Curricula. *Journal of Educational Multimedia and Hypermedia*. 2009;18(4):429–51.
- [5] Schwarz D, Komenda M, Dušek L, Šustr R, Šnábl I. MEFANET after four years of progressing: 4-D model for digital content quality assessment. Brno; 2010. 9. Harden RM. Looking back to the future: a message for a new generation of medical educators. *Medical Education*. 2011;45(8):777–84.
- [6] Masaryk university. OPTIMED - project documentation [Internet]. 2012. Available from: <http://opti.med.muni.cz/>.
- [7] Fact Sheet Medical Subject Headings (MeSH®) [Internet]. [cited 2012 Jul 12]. Available from: <http://www.nlm.nih.gov/pubs/factsheets/mesh.html>.
- [8] Needleman M. The Shibboleth Authentication/Authorization System. *Serials Review*. 2004;30(3):252–3.
- [9] Schwarz D, Dušek L. The MEFANET project [Internet]. [cited 2011 Mar 18]. Available from: <http://www.mefanet.cz/index-en.php>.
- [10] Šmerk P. K morfologické desambiguaci češtiny [Internet]. 2008 [cited 2012 Jul 15]. Available from: https://is.muni.cz/auth/th/3880/fi_r/.

- [11] Brandejs M, Hollanová I, Misáková M, Pazdziora J. In-house developed uis for traditional university: Recommendations and warnings. Proceedings of the 7th International Conference of European University Information Systems (EUNIS) [Internet]. 2001 [cited 2012 Jul 15]. p. 234–7. Available from: http://subs.emis.de/LNI/Proceedings/Proceedings13/49_In-houseDevelUISforTradUni.pdf.
- [12] Czech academic identity federation eduID.cz [eduID.cz] [Internet]. [cited 2011 Mar 18]. Available from: <http://www.eduid.cz/wiki/en/eduid/index>.
- [13] Šmerk P. Fast Morphological Analysis of Czech. RASLAN 2009 Recent Advances in Slavonic Natural Language Processing. 2009;13.

Verification of TaToo tools from the perspective of Validation Scenarios Solved by Masaryk University Team

Miroslav Kubásek, Jiří Hřebíček

Institute of Biostatistics and Analyses, Masaryk University
Kamenice 126/3, 625 00, Brno, Czech Republic
{kubasek, hrebicek}@iba.muni.cz

Abstract

The synthesis of existing Persistent Organic Pollutants pollution monitoring databases with epidemiological data is considered for identifying some impacts of Persistent Organic Pollutants on human health. This task requires new, rich, data, services and models discovery capabilities from a multitude of monitoring networks and web resources. The FP7 project TaToo (Tagging Tool based on a Semantic Discovery Framework) is setting up a semantic web solution to close the discovery gap that prevents a full and easy access to web resources. The use of TaToo tools together with the Global Environmental Assessment and Information System and the System for Visualizing of Oncological Data is discussed as TaToo validation scenario for anthropogenic impact and global climate change influence on Persistent Organic Pollutants trajectory.

Abstrakt

Syntéza současných databází obsahujících data z monitoringu perzistentních organických polutantů ve spojení s epidemiologickými daty je stěžejní pro identifikaci možného dopadu těchto látek na lidské zdraví. Tento úkol však vyžaduje mnohem bohatší data, nové služby a modely, včetně pokročilých přístupů k získávání informací o dalších dostupných datových zdrojích. FP7 projekt TaToo má za cíl vytvořit sémantický webový systém, který usnadní přístup k webovým zdrojům. V tomto článku popisujeme použití TaToo nástrojů spolu se systémy GENASIS a SVOD, které mimo jiné tvoří jeden z validačních scénářů tohoto projektu.

Key words

POPs, TaToo, SVOD, GENASIS, ontology

Klíčová slova

POPs, TaToo, SVOD, GENASIS, ontologie

1. Introduction

Persistent organic pollutants (POPs) represent a long-term problem which is connected with the production, application, and disposal of many hazardous chemicals and their impacts on human health. The Research Centre for Toxic Compounds in the Environment (RECETOX) of the Masaryk University (MU) is focused on the research of the fate and biological effects of POPs and other toxic substances in the environment. RECETOX monitors these chemicals in air, soil, water or human milk, and supports the implementation of international conventions on chemical substances like the *Stockholm Convention on Persistent Organic Pollutants*⁶.

RECETOX closely cooperates with the Institute of Biostatistics and Analyses (IBA) of MU. IBA is a research institute oriented to the solution of scientific projects and providing related services, especially in the field of environmental, biological and clinical data analysis. IBA created the *System for Visualizing of Oncological Data* (SVOD)⁷ – web portal of epidemiology of malignant tumours in the Czech Republic, which is based on the data from the Czech National Oncology Register [1].

⁶ Stockholm Convention on Persistent Organic Pollutants, <http://chm.pops.int/>

⁷ <http://www.svod.cz> - System for Visualizing of Oncological Data, <http://www.svod.cz/?sec=aktuality&lang=en>

Specific effects of POPs can include cancer, allergies and hypersensitivity, damage to the central and peripheral nervous systems, reproductive disorders, and disruption of the immune system. Some POPs are also considered to be endocrine disrupters, which, by altering the hormonal system, can damage the reproductive and immune systems of exposed individuals as well as their offspring; they can also have developmental and carcinogenic effects.

In January 2010 RECETOX launched the first version of the *Global Environmental Assessment and Information System* (GENASIS)⁸ – web portal which provides information support for implementation of the Stockholm Convention at international level. Initial phase of the GENASIS is focused on data from regular POPs monitoring programmes, providing a general overview of spatial patterns and temporal trends of pollutants concentrations. The aim is now to try to find out whether there is a connection between the concentration of POPs and cancer occurrence in some regions [4].

This task requires new discovery information and communication technology (ICT) tools which will be developed within the FP7 project TaToo⁹ (Tagging Tool based on a Semantic Discovery Framework) and shared the vision of a *Single Information Space in Europe for the Environment* (SISE) [7]. It aims to develop tools allowing third parties to easily discover web resources (data, services and models) and to add valuable information on to these resources. TaToo tools will be validated in three different validation scenarios. MU is solving the TaToo validation scenario of *Anthropogenic impact and global climate change*. It aims to improve the discovery of web resources in the domains of environmental pollution by POPs including influence of global climate change and epidemiology and tries to find relationships between these domains [2,3].

2. Scenario description

The MU scenario *Anthropogenic impact and global climate change* is dealing with the correlation of environmental pollutants and their health impact on the population and the correlation of transport of environmental pollutants with global climate change. The aim is to create a central place for researchers, domain experts and decision makers to discover and access interdisciplinary knowledge in more efficient and usable way that is the currently state of the art. Due to the fact that there is an enormous amount of information resources in scientific fields, which is steadily growing, available search mechanisms like search engines, scientific networks and similar technologies are not sufficient to meet the complex requirements of today's researchers and scientists. The result of conventional discovery processes are often not matching the domain context of the users and obligate them the tedious task of filtering large result sets to obtain the original object of the interest of the researcher intended to find with the search. Therefore the need arises for an improving discovery method, which will incorporate the domain knowledge and additional semantic information into the search in order to obtain a more fitting result for the specific context of the user.

The MU scenario not only aims to improve the discovery of scientific resources for one particular domain, but also tries to discover and create new relationships among different domains. The correlation of environmental pollutants including their transport due to global climate change and their health impact on the population is only one significant example of creating new relationships among different domains. These dependencies could represent new scientific insights for already available resources and connect the knowledge of the single domains. These relationships should facilitate further discovery process to deliver matching resources of multiple domains.

The MU scenario represents the close cooperation and joint venture of two university institutes: RECETOX and IBA .

⁸ <http://www.genasis.cz> - Global Environmental Assessment and Information System, <http://www.genasis.cz/main-index/en/>

⁹ Tagging Tool based on a Semantic Discovery Framework, <http://www.tatoo-fp7.eu/tatooweb/>



Figure 2: GENASIS analytical tool

RECETOX is an independent institute of the MU. RECETOX performs research, development, education and expertise in the field of environmental contamination by toxic compounds with specific focus on persistent organic pollutants (POPs), polar organic compounds, toxic metals and their species and natural toxins - cyanotoxins. It is also Stockholm Convention Regional centre for capacity building and transfer of technology in Central and Eastern European countries. The Stockholm Convention on Persistent Organic Pollutants (Stockholm convention) is a global treaty to protect human health and the environment from chemicals that remain intact in the environment for long periods, become widely distributed geographically and accumulate in the fatty tissue of humans and wildlife [8]. RECETOX is formed by several research divisions, service laboratories and technology-transfer centres: Environmental chemistry and modelling, Ecotoxicology and risk assessment, Trace laboratory, and Laboratory of data analyses. Research and development of the centre include monitoring of environmental matrices, studies of environmental fate and effects (ecotoxicology) of toxic compounds, ecological and human risk assessment as well as the development of informational and expert systems.

RECETOX project GENASIS provides information support for implementation of the Stockholm convention at an international level. The initial phase of the GENASIS project is focused on data from regular monitoring programmes of POPs, providing a general overview of spatial patterns and temporal trends of pollutants concentrations.

IBA is a research institute of the MU, which is focused on delivering solutions to research problems arising in the environment and human health and it is providing related services, especially in the field of biological and clinical data analysis, organization and management of clinical trials, software development and Information and Communication Technology (ICT) applications. IBA research activities are primarily focused on organizational and expert services for large scientific projects. IBA is formed by four divisions: Division of Data Analysis, Division of Clinical Trials, Division of Information and Communication Technologies, and Division of Environmental Informatics and Modelling. For example, IBA created SVOD - the first web portal for epidemiology of malignant tumours in the Czech Republic, system for visualizing of oncological data based on the data from the Czech National Oncology Registry.

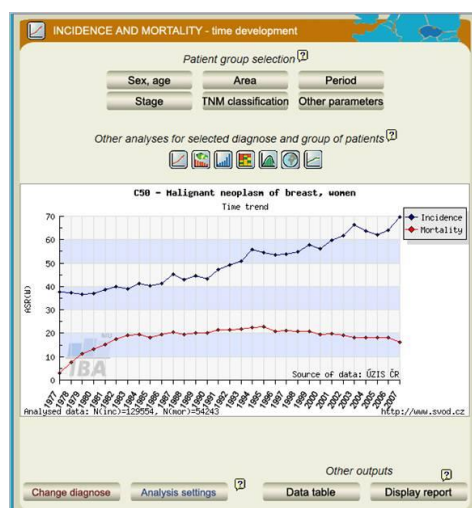


Figure 3: SVOD analytical tool

The objective of the MU scenario is to use and validate the resulting tagging and discovery framework of the TaToo project. Since the primary scope of the TaToo project is to facilitate the discovery of environmental resources, this scenario delivers the perfect opportunity to validate the resulting solution against challenging real word problems. There are numerous scientific domains available and actively researched at the MU, but two important domains have been carefully chosen to demonstrate and validate the envisioned functionality of the TaToo project. The vision of the MU scenario is that other scientific domains could follow the initial institutes to further spin a new kind of knowledge network to deliver a new generation of tools and methods to effectively and conveniently support the scientific user in their daily work.

3. Validation

Current TaToo development at MU is demonstrated in this section. To achieve the goal of the MU Scenario, it is necessary to use data from national and international monitoring networks, and to discover and obtain as-complete-as-possible data sets representing anthropogenic impact. There is a lot of various environmental data on the internet. However, by using current search engine it is difficult to find (all of) them and choose the relevant results only. These results are simpler to use in a real application. Discovery, use, and reuse of these data require enhancements of meta-information descriptions, which can be achieved through TaToo's semantic rich environment. In this context, MU intends to employ TaToo tools and will validate their performance for tagging and semantic rich discovery of resources of anthropogenic impact and the influence of global climate change on the transport of pollutants [5,6].

3.1 Validation resources

We categorized possible web resources for the our scenario into these categories:

- **Primary resources:** Structured raw data e.g. cancer patient records (diagnosis, sex, age, etc.) or measurements like time series of POPs (method, compound, substance etc.).
- **Secondary resources:** Aggregated or processed information based on primary data e.g. diagrams, analysis results, automatically generated reports, scientific publications, books etc. in form of well known datatypes (PDF, doc, txt etc.)
- **Information services:** Internet based services which provide information from the first and second category. For example Sensor Observation Services which provide Time Series for POPs in the form of compound measurement values (Web Services standard stc.).
- **Models:** Meta-information about mathematical and computational model for calculation of POPs distribution in the environment or dispersion models, which are used to estimate or to

predict (forecast) concentrations of airborne pollutants emitted from sources such as industrial facilities, local heating or traffic.

3.2 Type of Users

This chapter introduces shortly the different types of users who will use TaToo Tools. The users are divided into three categories:

- **Scientific users:** scientific users are regular users with scientific background and assumed IT skills. They will use the system to discover resources from both domains (POPs, health issues). They will be able to find resources, find similar ones (having already found some resource), compare the resources, and also to find connections between resources. Everything on the “read only” basis.
- **Domain experts:** group of domain experts collects users who have some additional functionality to scientific users. Domain experts can also evaluate resources and assign metadata to the resources. By the means of mentioned functions they will contribute to the information enrichment process.
- **Administrators:** administrators will be responsible for organisational and maintenance tasks in order to guarantee proper system functionality. This involves also user administration, system settings, problem solving, user support etc.

3.3 Proposed Use Cases

Proposed Use Cases of Validation Scenario 3 are as follows:

- **Use Case 1:** Discover resources with existing tools - SVOD and GENASIS users will be provided with the possibility to indirectly use the TaToo functionality for discovering similar resources based on analysed objects in mentioned web portals.
- **Use Case 2:** Generic discovery - the goal of this use case is to deliver improvements regarding result relevance compared to conventional search engine results. The relevance of the resources to the search criteria should be improved so that the user receives more potential interesting search results. This circumstance hopefully leads to a reduction in the tedious effort of scanning the search results for matching entries.
- **Use Case 3:** Persistent Organic Pollutant Resource Discovery.
- **Use Case 4:** Oncological resource discovery. Use Case 3 and Use Case 4 bring extended domain specific search. Domain experts are allowed to enrich the resources by using TaToo Tools. This example of use cases will be in detail describe in the next chapter.
- **Use Case 5:** Define resource uncertainty - domain experts are allowed to define certain quality criteria for resources.
- **Use Case 6:** Compare resources - enables users to compare found resources on the fly after the discovery. This Use Case 6 should be helpful in finding the connections between different resources either in the same domain or in different domains.
- **Use Case 7:** Find similar resources. The Use Case 7 brings the functionality to search for similar resources based on interesting resource already found.
- **Use Case 8:** Find related resources. The last Use Case 8 provides with searching for related resources in other knowledge domains based on an already found resource.

4. Validation Case Studies

The purpose of this chapter is to introduce validation case studies to be used to validate TaToo Tools. For each use case we provide a short description and explain their purpose.

4.1 Case Study 1: Discover resources with existing tools

Purpose of this case study is to provide the users of SVOD and GENASIS portals with the possibility to indirectly use the TaToo functionality for the discovery of similar resources based on analysed objects. The TaToo discovery is started directly from within the web analysis tools. The relevant

information needed for the search would be already entered via the web interface during the analysis and can be submitted to the TaToo framework.

In this case study we implemented simply "TaToo button" (see Figure 2 and Figure 3) into SVOD analyse tools and applied TaToo Discovery Service to retrieve a list of related data resources from TaToo Repository (see Figure 4). Queries used in calling the Discovery Service are automatically built from actual settings of SVOD analytics tool. Communication between TaToo Discovery Service and developed discovery application is based on SOAP standard (see Figure 1). Accordingly, Case Study 1 will cover Use Case 1 and will use Discovery service.

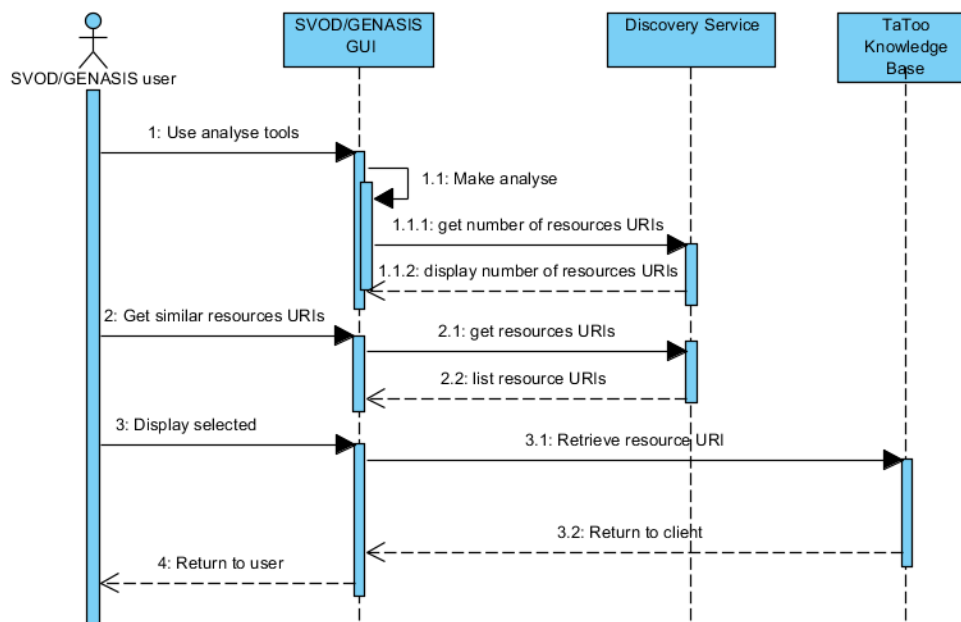


Figure 1. Sequence diagram of discovery related resources

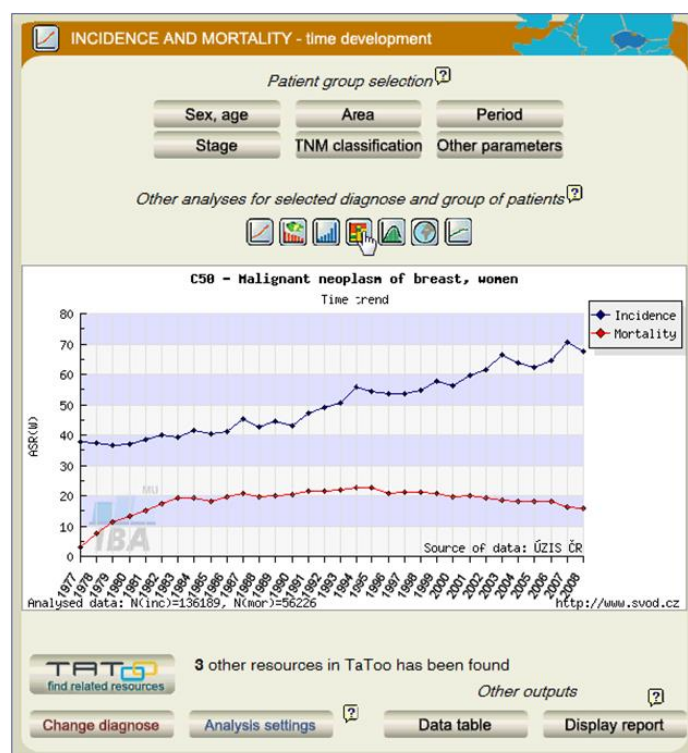


Figure 2. Integration of TaToo functionality into SVOD analysis tool

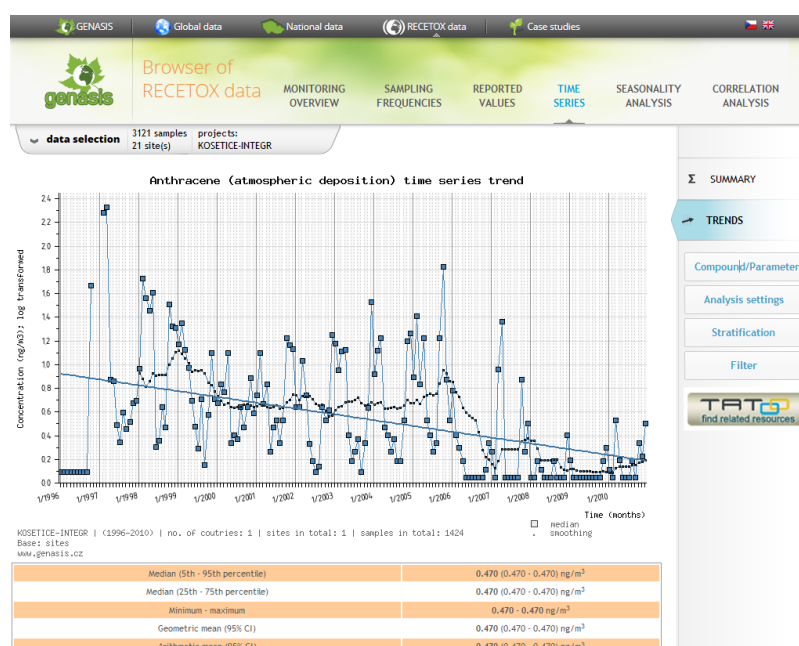


Figure 3. Integration of TaToo functionality into GENASIS analysis tool

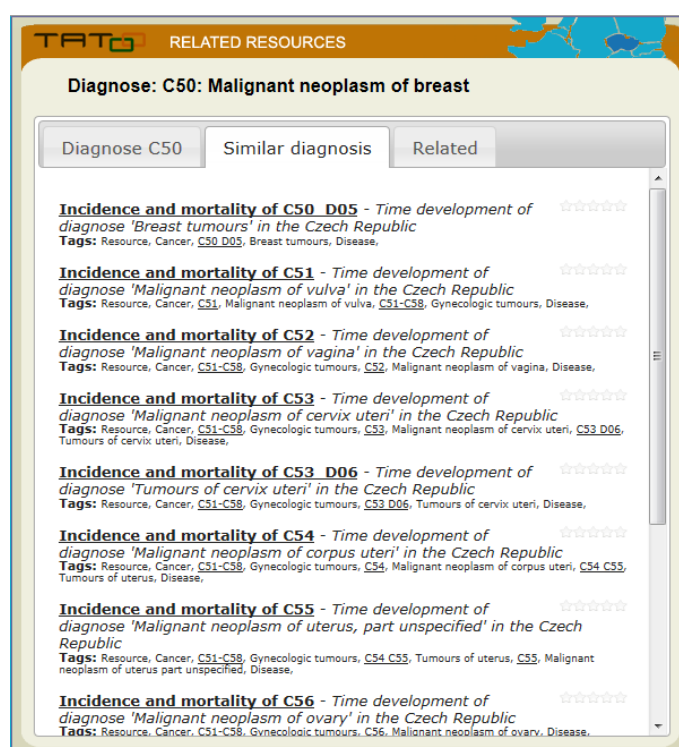


Figure 4. List of similar resources in SVOD portal (live demo)

4.2 Case Study 2: Generic discovery

The user wants to discover resources of a particular domain of interest matching certain criteria and keywords. The goal of this use case is to deliver improvements regarding result relevance compared to conventional search engine results (see Figure 5 for user interface of these case study). The user wants

to find from the multitude of available resources the most interesting for his particular case, which is represented by the input information. The found result should therefore have a higher probability to fit the desired domain context.

Additional to domain specific information the system should also include other dimensions in the discovery like time range and geospatial information of the results, in order to further specify the domain of interest. The resulting list of resources should include additional information to the resource such as relevance to the search query, uncertainty information about the quality of the resource contains, file type of the resource etc. The search should not only deliver results in the original language used to specify the search query, it should also deliver results in foreign languages which match the domain context (e.g. user type search query in Czech language and TaToo Tools will be able to understand and give to the user also English resources fulfilling typed query).

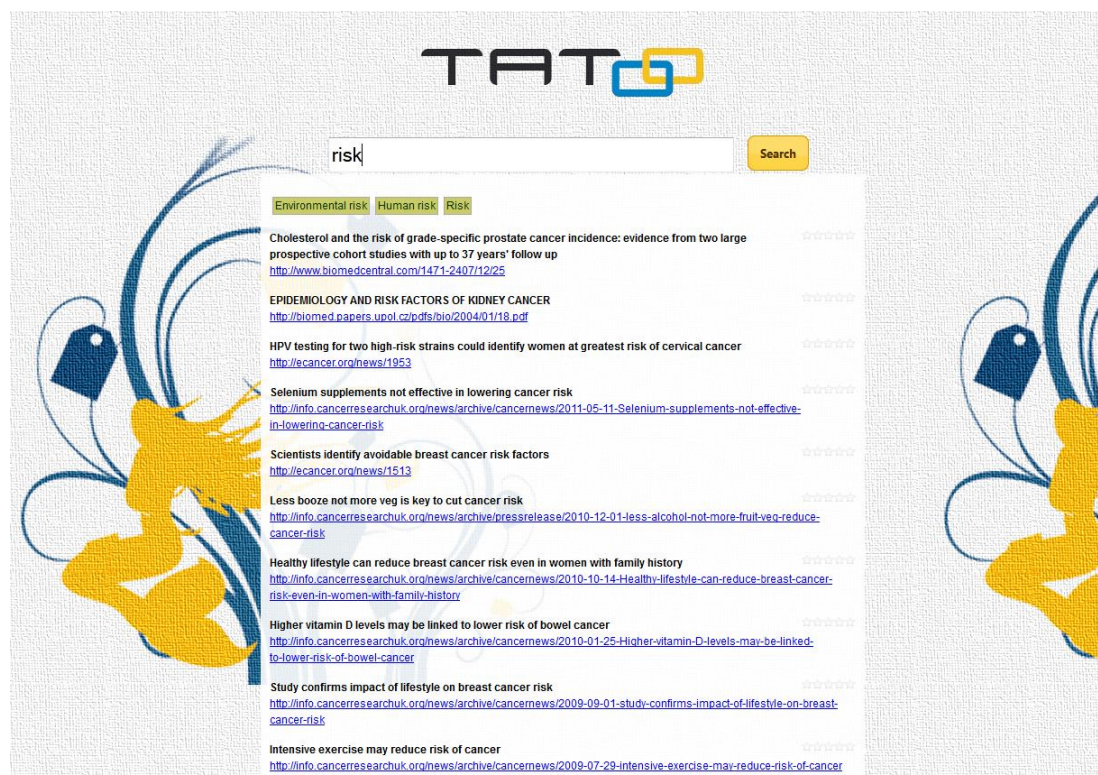


Figure 5. Interface for TaToo generic discovery

This Case Study 2 will cover Use Case 2 and is interpreted by Simple Search Portlet. The first version of this implemented portlet contains only a part of desired functionality, so only simple queries can be processed.

4.3 Case Study 3: Specific Resource Discovery

A user wants to find information in the domain of bio chemistry, specifically about persistent organic pollutants monitoring. The interesting resources range from primary data such as time series with the actual measurements and additional information about measurement methods, measured compounds, etc. to high level information are generated from this primary data such as statistics and time trends of pollutants.

Similar to biochemistry domain (see Figure 6) specific resource search we can discover resources with the focus on oncology domain (see Figure 7). The user is interested in the discovery of cancer related resources. Nonetheless researchers are interested in a discovery of resources containing evaluations statistics, and reports regarding cancer incidences and mortality rates. Similar to persistent organic pollutant resource discovery the user wants to discover resources based on a domain specific search

mask with common parameter such as diagnosis, gender, patient number, etc. The case study 3 will cover Use Cases 3 and 4.

POPs Discovery

Code of chemical: Filter

Name of chemical: Number of sites:

Matrix: Number of samples:

Fraction:

Sampling method:

Keywords:

Timerange

Start Year:

End Year:

Geospatial Region **Results** **Similarity Search** **Relationship Search**

☐ by Text

Country:

Region:

City:

☒ by Geometry

Custom Geometry

Map:

Figure 6. POP discovery mock-up.

Cancer Discovery

Diagnosis: Interest

Name: ☒ Mortality

Sex: ☒ Incidence

Age: Mortality/Incidence

Filter

Patient Count:

Computing method:

Timerange

Start Year:

End Year:

Geospatial Region **Results** **Similarity Search** **Relationship Search**

☐ by Text

Country:

Region:

City:

☒ by Geometry

Custom Geometry

Map:

Figure 7. Cancer discovery mock-up

4.4 Case Study 4: Find similar and related resources

This case study represents the functionality of the search for similar (related) resources based on already found interested resource. If the user finds a resource that matches his needs a new search is started based on a current resource. The search for related resources means found resources in other

knowledge domains based on an already found resource. For example the user wants to find pollutant monitoring data for a specific time range and geospatial region, based on the values of a discovered cancer analysis. The geospatial extend and temporal extend from the cancer analysis will be used to perform a new search. The user only has to provide and specify the domain of interest in which new resources should be discovered.

This complex Case Study 4 will cover Use Cases 1, 2, 3, 4, 7, and 8.

4.5 Case Study 5: Tagging multiple SVOD/GENASIS resources

The main goal of this validation use case is to support for SVOD/GENASIS resources to publishing and add them into the TaToo Knowledge Base. The RDF resource descriptions are generated in accordance to the TaToo specifications of the Minimal Environmental Resource Model (MERM). This use case will use proposed domain ontology (see The MU Ontology Development – Internal Report)[9] to test TaToo System. Generated RDF descriptions of resources will be stored as XML files on server or can be directly by Tagging Web Services stored into the TaToo Repository.

For simplification of harvesting process MU provided "Resource Catalogue"¹⁰ which contains all possible resources from MU domains (both POPs and cancer). This catalogue enable to our administrators to manage resources (edit basic resource attributes and keep them up-to date) and also manage topics from our domain ontology connected to these resources. This catalogue also provides RDF connectors to all resources to enable harvesting process. In *Figure 8* you can see the screenshot of this tool.

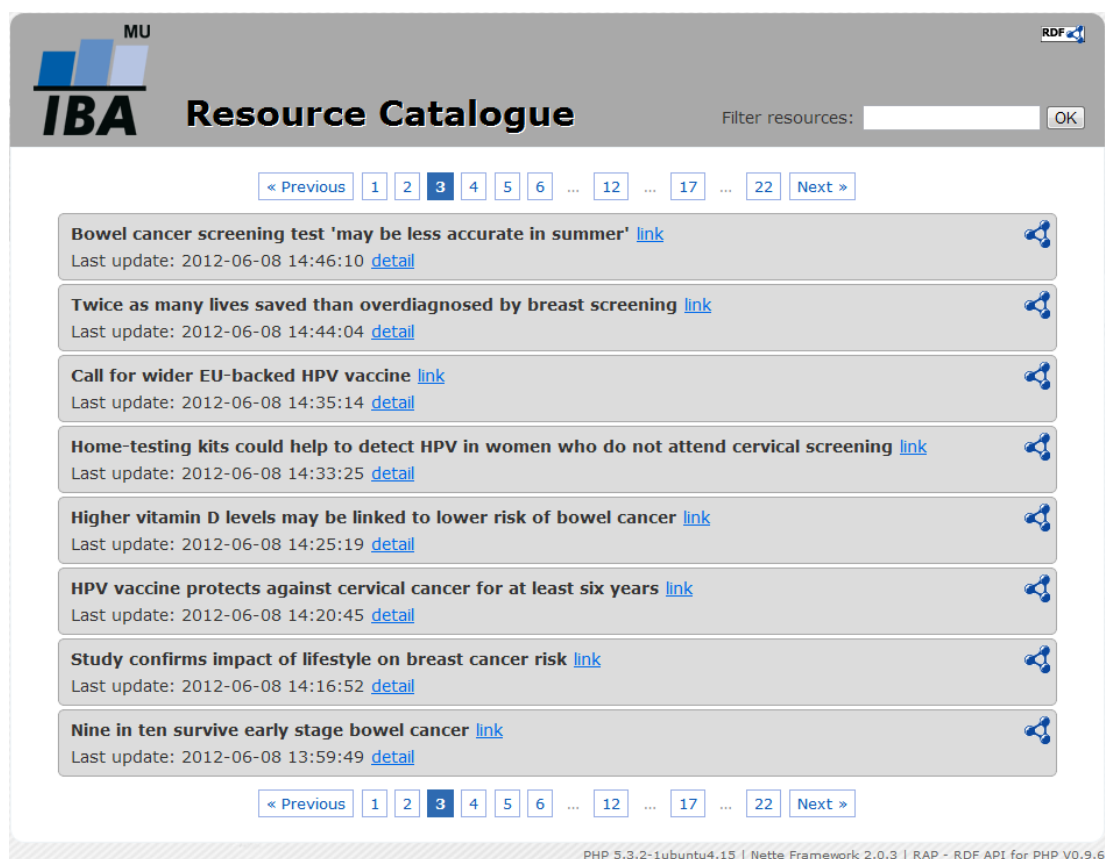


Figure 8: Resource Catalogue (<http://ontology.genasis.cz>)

¹⁰ Available on <http://ontology.genasis.cz>

4.6 Case Study 6: Tagging a specific data resource

This Case Study allows a regular user of TaToo Tools to tag information about a discovered or already known resource, where a resource could be a data source, a service, a Web page, etc. The user is prompted by the portlet with a selection panel to choose terms from an ontology to create tags for a resource or a set of resources. When the user adds tags, the Tagging Portlet contacts the TaToo Tagging Service to update the information related to the resource.

This case study is realized by using several implemented portlets (e.g. Tagging Generic Portlet and MU Specialized Tagging Portlet).

5. Conclusion

This contribution describes how the anthropogenic impact and influence of global climate change scenario has been implemented taking advantage of several TaToo Components.

The validation of TaToo Tools in the context of the Validation Scenario 3 is realised by means of six validation case studies covering several ways how to discover desired data resources from the TaToo Repository. The case studies also offer possibility to enrich TaToo Repository with new data by appropriate tagging of discovered resources.

The overall implementation and integration of TaToo into the SVOD and GENASIS portals have made good progress. A few issues remain which have to be done in the final version of TaToo:

- Multilingualism: A new domain ontology with multilingual labels has been created, and has to be integrated in the upcoming final portal.
- TaToo Linked Data approach: adding relationship links between TaToo resources according to Linked Data principles. We expect that this feature will enable our Validation Scenario for multi-domain search.

6. Acknowledgements

“The research leading to these results has received funding from the European Community's Seventh Framework Programme (FP7/2007-2013) under Grant Agreement Number 247893.”

7. Bibliography

- [1] Dušek, L., Mužík, J., Koptíková, J., Brabec, P., Žaloudík, J., Vyzula, R., Kubásek, M.: The national web portal for cancer epidemiology in the Czech Republic. In: *Enviroinfo 2005. 19th International Conference Informatics for Environmental Protection*. pp. 434--439. Masaryk University Press, Brno ISBN 80-210-3780-6 (2005)
- [2] Hřebíček, J., Dušek, L., Kubásek, M., Jarkovský, J., Brabec, K., Holoubek, I., Kohút, L., Urbánek, J.: *Anthropogenic Impact and Global Climate Change. Description of Validation Scenario in TaToo Project*. In Hřebíček J., Pitner T., Ministr J.. 7. letní škola aplikované informatiky. Indikátory účinnosti EMS podle odvětví. Brno: pp. 6-23, nakladatelství Littera, ISBN 978-80-85763-59-1. (2010)
- [3] Hřebíček, J., Dušek, L., Kubásek, M., Jarkovský, J., Brabec, K., Holoubek, I., Kohút, L., Urbánek, J.: *Validation Scenario for Anthropogenic Impact and Global Climate Change for TaToo*. In *Proceedings of the Workshop "Environmental Information Systems and Services - Infrastructures and Platforms"*. CEUR-WS, Aachen. ISSN 1613-0073 (2010)
- [4] Klánová, J., Čupr, P., Holoubek, I., Borůvková, J., Příbylová, P., Kareš, R., Kohoutek, J., Dvorská, A., Komprda, J.: *Towards the Global Monitoring of POPs - Contribution of the MONET Networks*. RECETOX, Masaryk University, Brno, ISBN 978-80-210-4853-9 (2009)
- [5] Kubásek, M., Hřebíček, J., Kalina, J., Dušek, L., Holoubek, I.: *Semantics Annotations of Ontology for Scenario: Anthropogenic Impact and Climate Change Issues*. In Jiří Hřebíček, Gerald Schimak, Ralf Denzer. In *Environmental Software Systems. Frameworks of*

- eEnvironment. 9th IFIP WG 5.11 International Symposium. ISESS 2011. Heidelberg: pp. 398-406, Springer, ISBN 978-3-642-22284-9 (2011)
- [6] Pariente, T. et al.: A Model for Semantic Annotation of Environmental Resources: The TaToo Semantic Framework. In Jiří Hřebíček, Gerald Schimak, Ralf Denzer. Environmental Software Systems. Frameworks of eEnvironment. 9th IFIP WG 5.11 International Symposium. ISESS 2011. Heidelberg : Springer. ISBN 978-3-642-22284-9 (2011)
 - [7] Rizzoli, A., Schimak, G., Donatelli, M., Hřebíček, J., Avellino, G., Mon, J.: TaToo: tagging environmental resources on the web by semantic annotations. In: iEMSS 2010. International Congress on Environmental Modelling and Software. Modelling for Environment's Sake. iEMSS, pp.1192-1199, Ottawa. ISBN 978-88-903574-1-1 (2010)
 - [8] Urbánek J, Brabec K, Dušek L, Holoubek I, Hřebíček J., Kubásek M.: Monitoring and Assessment of Environmental Impact by Persistent Organic Pollutants. In: Diamantaras K, Duch W, Iliadis L, editors. Artificial Neural Networks -- ICANN 2010. Vol 6354. pp. 483-488, Springer, Heidelberg, ISBN 978-3-642-15824-7 (2010)
 - [9] Avellino G., Kubásek M., Hřebíček J. : The MU ontology development, TaToo Project, internal report, (2011)

Tagging Tool based on a Semantic Discovery Framework - Semantic Framework Implementation

Miroslav Kubásek¹, Jiří Hřebíček¹, Sinan Yurtsever², Pascal Dihé³, Sasa Nesic⁴, Giuseppe Avellino⁵, Luca Petronzio⁵

¹Institute of Biostatistics and Analyses, Masaryk University,
Kamenice 126/3, 625 00, Brno, Czech Republic

²ATOS Origin, Madrid, Spain

³cismet GmbH, Altenkesseler Strasse 17 D2, 66115 Saarbrücken, Germany

⁴IDSIA, Lugano, Switzerland

⁵Telespazio S.p.A, Via Tiburtina, 965 - 00156 Rome – Italy
{kubasek, hrebicek}@iba.muni.cz

Abstract

The present document has been produced by the consortium of the European Project FP7-247893 Tagging Tool based on a Semantic Discovery Framework (TaToo). It describes the third and final iteration of the implementation of the TaToo Semantic Service Environment and Framework. This document reports on the implementation activities carried out in tasks T4.1, T4.2 and T4.3, acting as glue of all the implementation tasks of WP4.

Abstrakt

Následující dokument je výsledkem práce týmu evropského projektu FP7-247893 Tagging Tool based on a Semantic Discovery Framework (TaToo). Popisuje třetí a finální verzi implementace prostředí sémantických služeb a rámec projektu TaToo. Tento dokument představuje popis implementace jenž je výsledkem prací teamu WP4.

Key words

TaToo, Ontology, Semantic web

Klíčová slova

TaToo, Ontologie, Semantický web

1. Introduction

The aim of the TaToo Framework is to provide an infrastructure to fill the discovery gap between environmental resources and users. The TaToo Framework enables experts as well as arbitrary users to share trusted and reliable information and to allow easy discovery and semantically enhanced tagging of existing environmental resources.

The major achievements and results of the final iteration of the TaToo Framework implementation are:

- The decision on implementation issues based on existing state-of-the-art solutions on the semantic discovery, tagging and Linked Data fields;
- The consolidation of the TaToo Framework Architecture implementation viewpoint;
- The realization of the TaToo Portal as a showcase client of the TaToo Public Services using an unified technology;
- The implementation of the final version of the main components of the TaToo Framework, including new components and a better, simple and usable version of the TaToo Portal;
- The provision of the final ontology framework allowing cross-domain search, Linked Data extension, similarity between resources and multilingual aspects, following deliverable D3.1.3 where the basis for the TaToo ontology framework design was established;
- A detailed description of all identified components. It follows deliverable D3.1.3 where the basis for the detailed design of the TaToo Framework was established;
- The software implementation (code release) of the TaToo Framework V3 is also part of D4.1.3. The appropriate licensing schema will be provided in the final exploitation

deliverable, but most of the work done within the scope of WP4 follows an Open Source license.

2. Framework design and implementation description

This chapter describes the infrastructure supporting each Building Block and the implementation of the TaToo components. Conceptually, the TaToo Framework is designed as n-tier architecture composed of two main High Level Building Blocks:

1. **TaToo Public Services:** The externally visible face of the TaToo Core (TaToo System Components). To achieve maximum interoperability, these services must be accessible over standardised web service interfaces.
2. **TaToo Core Components:** Which provide the business logic, data and meta-information management. These components are only accessed by other System Components and by the TaToo Public Services. Consequently, interoperability is less of a concern and the System Components need not to be accessible over a standardised web service interface.

Figure 1 provides a complete overview on the Functional Purview of the TaToo Framework Architecture as defined in TaToo deliverable D313 [1].

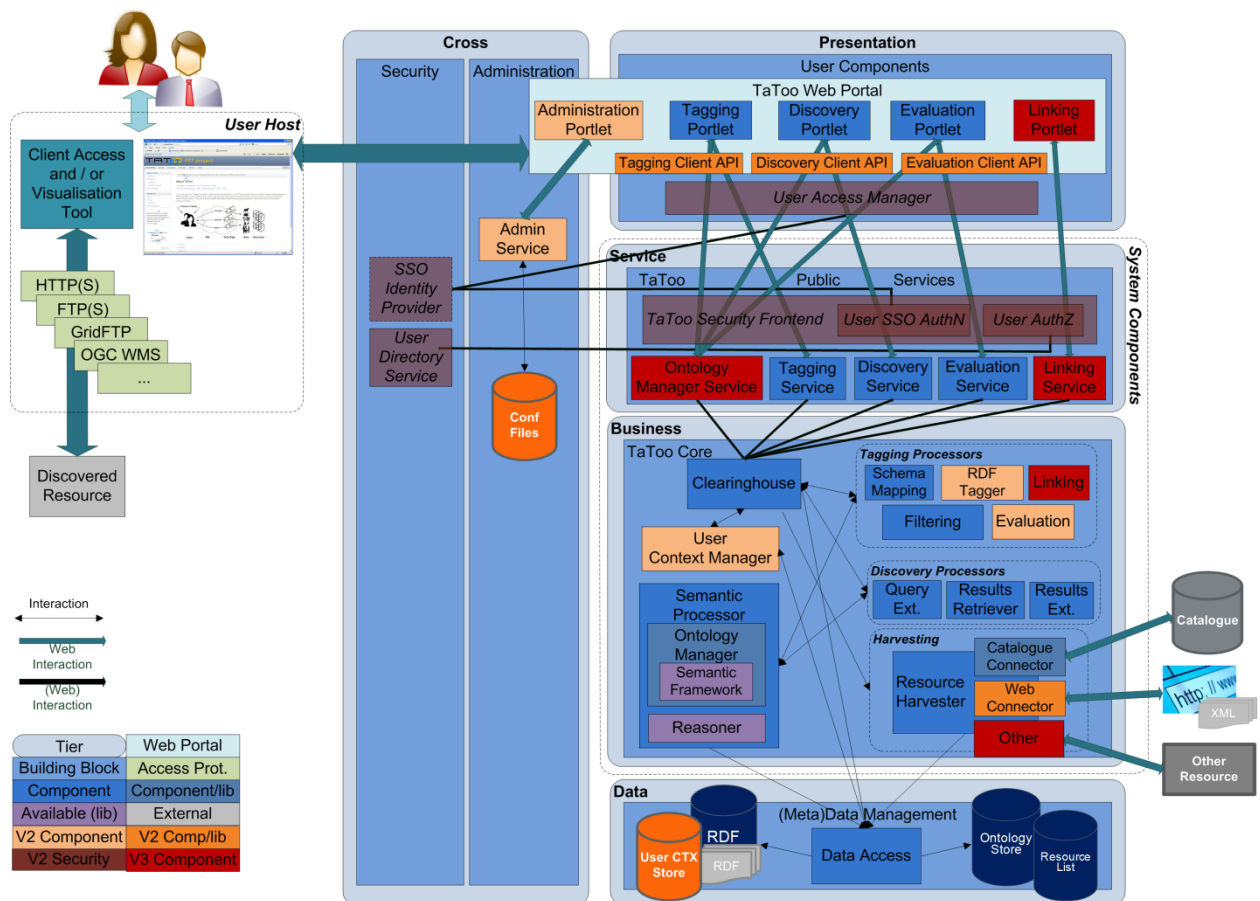


Figure 1: TaToo Framework Architecture v3

Please note that components of the data tier, although conceptually belonging to a Building Block, in general, with the exception of the Data Access Component, they represent data storage solutions and products like Semantic Repositories, RDBMS, etc. and thus do not need a dedicated description or specification.

2.1 TaToo Framework implementation view

This section defines the mapping between the Functional purview of the architecture drawn in D3.1.3 and the implementation decided in this deliverable.

2.1.1 Presentation tier

The Presentation tier contains all the components the user can take advantage of to access the functionality provided by the TaToo Framework (in particular the tagging and discovery functionality). The tier has been designed as made of a single Building Block, the User Components Building Block, containing a set of components to be exploited by the end user. These components can be of different type, such as portals (made of a set of portlets), (rich-) client application, browser plug-ins and any other type of tool¹¹. To access the TaToo Framework, they normally act as Web clients of the TaToo Web services (see Service tier in the next section).

For the final implementation phase, TaToo continued the implementation of a Portal, see TaToo deliverable D313 [2]. The TaToo Portal is made of a set of portlets, which can be seen as tools as well, implemented basing on JSR168 and JSR286. Relying on the aforementioned JSRs is important while integrating the portal portlets coming from different technical partners.

The implementation of a portal as a first entry point and showcase to access the TaToo Framework has been motivated by the TaToo vision of providing easy access to users, thus only requiring a browser and an Internet connection wherever the users are located (no need to install other software or plug-ins). The portal personalisation aspect is also an important factor as the user can work comfortably using their personalised set of portlets, configured considering the required functionality, comfortably remaining bounded to their well known application domain. For the design and implementation description of the individual components (portlets) and their mapping to specific implementation technologies, please refer to section 2.1.1.

Currently the following portlets can be identified:

- The Tagging Portlets: Tagging, Geo-Tagging, Simple Tagging, Tags Editing;
- The Discovery portlets: Hierarchical Search, Simple Search, SPARQL, and Result Presentation portlets;
- The Evaluation Portlets: Add Evaluation and Evaluation Browser Portlets;
- The Linking Portlets: Linking and Links Browser Portlets (new in this iteration);
- The Administration Portlet.

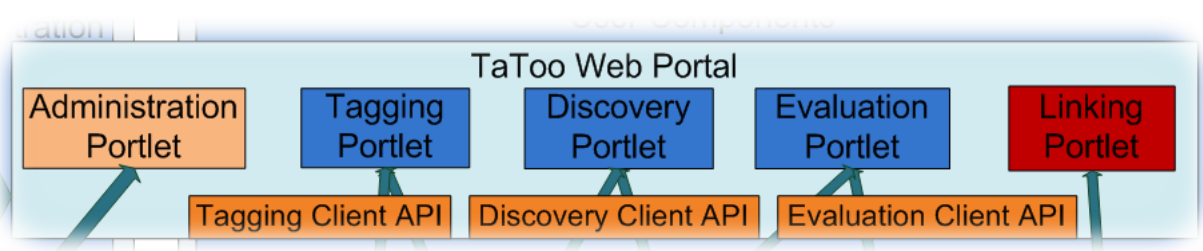


Figure 2: TaToo Web Portal

The reference implementation of the TaToo portal is based on Liferay Community Edition (Liferay) technology (providing the Portlet Container in Figure 3). Liferay provides the possibility to choose among a large set of alternative Portal Servers (bundles), such as Tomcat, Glassfish, JBoss, and others. Both Tomcat and Glassfish bundles are being evaluated mainly focusing on respective performance.

¹¹ In the context of TaToo, the term 'tool' is intended as a front-end component, generally with a Graphical User Interface (GUI), which allows the user (residing in the Presentation tier) to interact with the system taking advantage of the provided functionality.

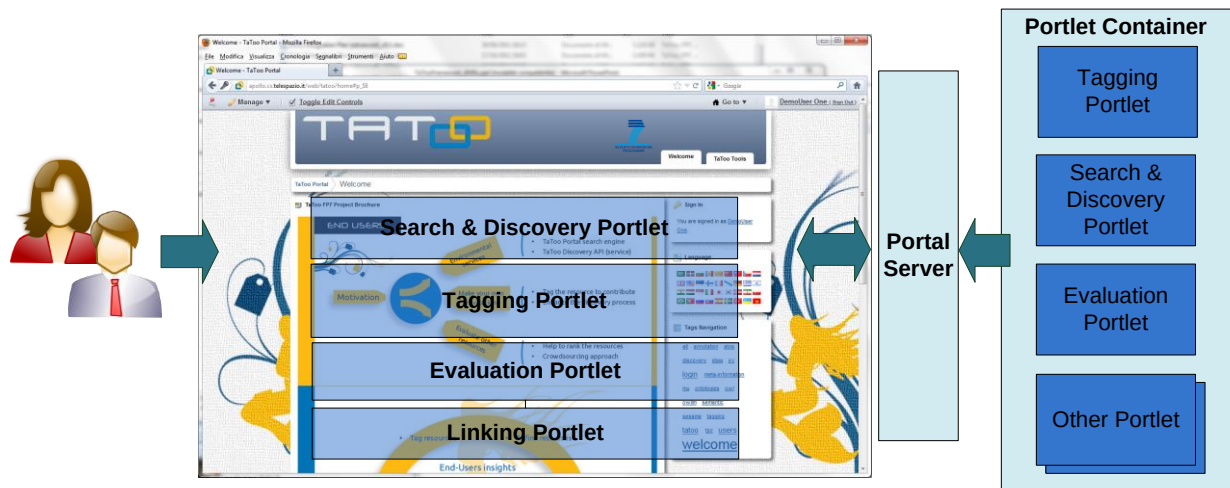


Figure 3: Portlets composing the TaToo Portal

In the final version of the TaToo Portal the following modifications have been adopted on the previous V2 version portlets to implement V3 portlets:

- Reduction of UI elements composing the portlets in order to provide simple use case workflows avoiding redundancies, i.e. TaToo resources and annotations are displayed only in the Result Presentation Portlet and have been removed from Tagging Portlet;
- Definition of new Portlet Events, that, together with existing ones, will cover each inter-portlet foreseen communication;
- Enable reading and displaying Ontology Individuals Labels instead of their URI;
- Minor graphic design updates to have a more appealing portal page.

This new approach allows the user to follow a simple and coherent workflow of actions when dealing with the TaToo Tools in the portal in order to take full advantage of TaToo functionalities. User Interface areas that were previously pertaining to a certain portlet to display redundant information are now shared together providing a harmonised and coherent TaToo Tools portal page.

2.1.2 Service tier

This chapter defines the mapping of the Service Tier (TaToo Public Services) to a web services platform which is described in accordance to the rules of the OASIS Reference Model for Service Oriented Architecture 1.0 [3]. A SOAP Web Services based approach is followed for the service specifications and implementations. For the design and implementation description of the individual services and their mapping to specific implementation technologies, please refer to section 2.1.8.

The decision to follow the SOA-RM for the specification of the basic properties of the TaToo SOAP Web Services Service Platform, and thus realising the TaToo Public Services as W3C compliant Web Services, is based on several general enterprise requirements TaToo deliverabl D233[4]. Namely, the most important requirements considered are:

- Use of concepts and standards
The usage of W3C standards decreases dependence on vendor-specific solutions and helps ensure the openness of TaToo Service Tier.
- Loosely coupled components
TaToo's service oriented architecture and especially the Service Tier as mediator between TaToo Tools and TaToo Core Components facilitate loose coupling.
- Extensibility / Flexibility
The Service Tier facilitates the integration of new services with additional functionality into a

TaToo semantic framework. Furthermore, TaToo Public Services offer the possibility to access TaToo functionality by new third party applications (custom clients).

The specification of the TaToo SOAP Web Services Service Platform comprises of specifications and descriptions of selected technologies required to formally specify the TaToo Public Services. Furthermore, it provides an informal description of the mapping from the functional specifications to implementation specifications. Please note, that the specification of the TaToo SOAP Web Services Service Platform does not impose any constraints on the individual service implementation, e.g. usage of a specific programming language, middleware, SOAP implementation, etc. Its primary objective is to ensure interoperability by enforcing that all service interfaces of the Service Tier are specified in the same manner and follow the same web service standards.

SANY FP6 project¹² recommends that the specification of a service platform shall be conformant to the OASIS Reference Model for Service Oriented Architecture 1.0 [5]. As a consequence, the TaToo SOAP Web Services Service Platform has to be described by a set of predefined platform properties.

Chyba! Nenalezen zdroj odkazů.: Lists the properties of the TaToo SOAP Web Services Service Platform

Property	Value
Platform Name <i>Name of the platform.</i>	TaToo SOAP Web Services Service Platform
Reference Model <i>Reference model on which the platform specification is based.</i>	W3C Web Services Architecture (W3C, 2004)
Interface Language <i>Formal machine-processable language used to define the service interfaces.</i>	Web Service Description Language (WSDL), Version 1.1 (W3C, 2001)
Execution Context <i>Information about preferred protocols, semantics, policies and other conditions and assumptions that describe how a service can and may be used, e.g. the specification of the transport and the security layer, the format of the messages exchanged, etc.</i>	Transport Protocol and Message Format: SOAP 1.2 HTTP binding as defined in SOAP Part 1: Message Framework, Version 1.2 (W3C, 2003) and Hypertext Transfer Protocol (HTTP), Version 1.1 (W3C, 2006).
Schema Language <i>Specification of the schema language used to define the information exchanged.</i>	General schema language: eXtensible Markup Language (XML) 1.0. Schema Language for semantic information models: RDF, RDFS and OWL, in particular a subset of OWL 2 called OWL2 RL.
Schema Mapping	Not applicable.

¹² <http://www.sany-ip.eu/>

Property	Value
<i>Specification of how to map the abstract level (e.g. UML) to the schema language used for this particular platform.</i>	
Information Model Constraints <i>Specification of the constraints on the Information Model, especially the constraints on the message format.</i>	TaToo ontologies shall be based on existing W3C standards, particularly RDF, RDFS and OWL, in particular a subset of OWL 2 called OWL2 RL.

The mapping of a functional specification of a TaToo Public Service to the correspondent formal implementation specification is carried out as a 1:1 mapping. This means that interfaces specified in a functional specification shall be mapped to WSDL portType elements, operations shall be mapped to WSDL message type elements, and parameters shall be mapped to corresponding WSDL message elements.

Since the parameters of an operation are rather generally and abstractly described in a functional specification, concrete types have to be specified in the implementation specification. The result of such a mapping is an implementation specification provided in the document that overwrites and complements the specification done in D313 [2]. Currently the following services can be identified:

- Tagging Services
- Discovery Services
- Evaluation / Validation Services
- Ontology Manager Services
- Similarity Services
- Linked Data Services

2.1.3 Business tier

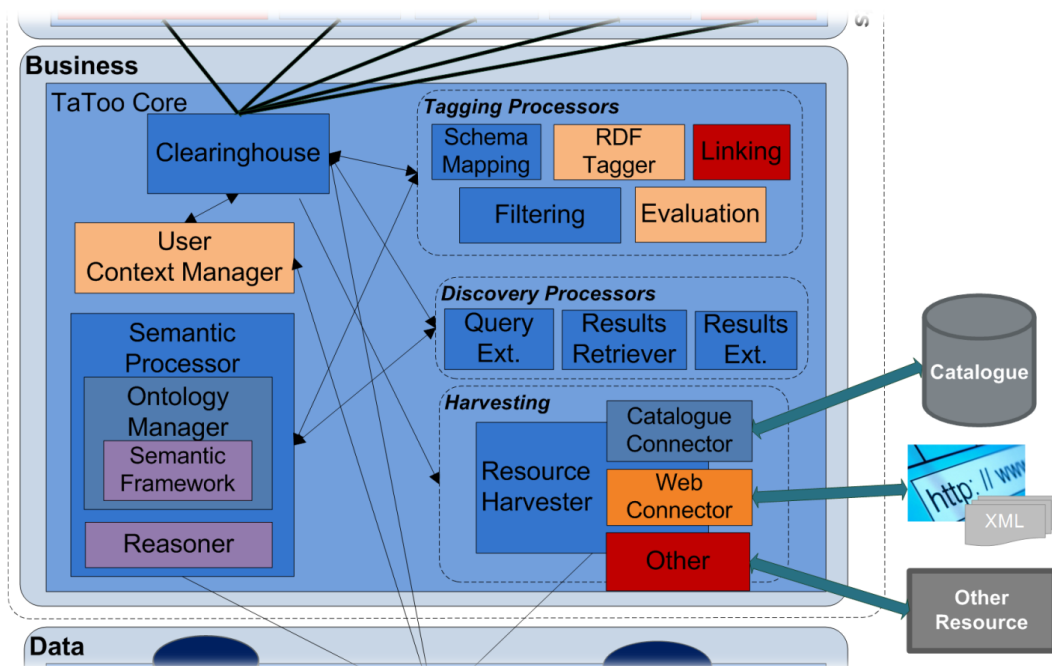


Figure 4: TaToo Core Components Building Block

As already stated, the Core Components Building Block provides the business logic of the TaToo System (see Figure 4). The main components of this Building Block are the:

- Clearinghouse: Service and entry point to the TaToo Core Components. The Clearinghouse provides a coherent interface between TaToo Public Services and Core Components. All communications happen through the Clearinghouse, which also contains the required business logic to orchestrate component calls (or invocations), where necessary;
- Semantic Processor: Its functionality is basically implemented using Sesame and OWLIM, covering the implementation and usage of:
 - Ontology Manager
 - Reasoner
- Tagging Processors: The Schema Mapping, the Visualisation and Filtering, the RDF Tagger, the Evaluation, and the Linking components;
- Discovery Processors: Reference implementation of the Query Expansion, Resource Retriever, and Resource Expansion components;
- Resource Harvester, together with a set of Resource Connectors:
 - JRC Resources catalogue connector;
 - Masaryk University Resources catalogue connector;
 - GENESI-DEC OGC Resources catalogue connector;
 - LinkedData Resources Connector;
 - Web Sites RDFa Connector.

2.1.4 Cross tiers

In the second phase of the TaToo Framework Architecture definition, new cross tiers have been added and defined in details: Security Cross tier and Administration Cross tier.

The Security and Administration Cross tier are depicted in the Figure 5 below.

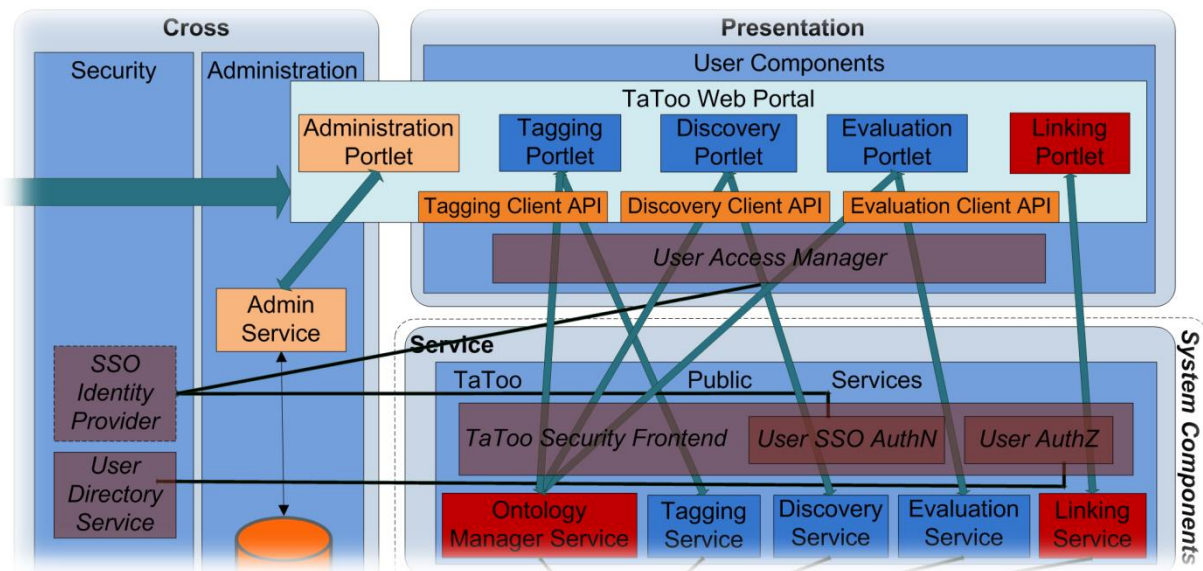


Figure 5: Security and Administration Cross tier

- Administration Service: A service to manage configurable preferences of components.
- Administration Portlet: Simple GUI of the Administration Service.
- TaToo Identity Provider: A service to manage Federation and user identities in the form of SAML Assertions. Based on the Shibboleth Identity Provider Service.

- **User Directory Service:** A service that manages user authorization based on its attributes, roles, groups, communities. The service is based on the LDAP Service.

2.1.5 Security

Considering the wide community that potentially will take advantage of the TaToo Tools and Framework, while designing the Security Cross tier the focus has been to realise a user friendly approach on the client side but robust enough on the Public Service tier side. Thus the TaToo Security architecture components realise a Single Sign-On infrastructure that allows the user to access all TaToo functionality, tagging, discovery, and evaluation, authenticating only once per session.

The brown boxes in Figure 5 are describing the components that are related to the Security infrastructure. In the presentation tier the client side authentication operation is realised through the User Access Manager, on server side the Public Service tier is secured by using a TaToo Security Frontend that acts as a secured proxy server toward Web Services operations calls, and finally in the Security Cross tier, the TaToo Identity Provider and the TaToo User Directory Service are indicated.

Concerning technologies used, TaToo Security components are based on Shibboleth (Shibboleth). Shibboleth is an open source Single Sign-On standard for web authentication across different organization, namely Federations, based on SAML (*Security Assertion Markup Language*) standard. It allows the service and data providers to make authorization decisions for access to the web protected resources based on the specific attributes of the user identity. Attributes are defined in Shibboleth with a standard format, and Federations can use both built-in attribute support and defined new ones. The Shibboleth system supports Federations by defining also formats for managing site configuration information and providing functionalities to create, distribute, and import that information. In TaToo, Shibboleth Federations concept is useful to distinguish between different type of users, depending on their domain or on their role in the community.

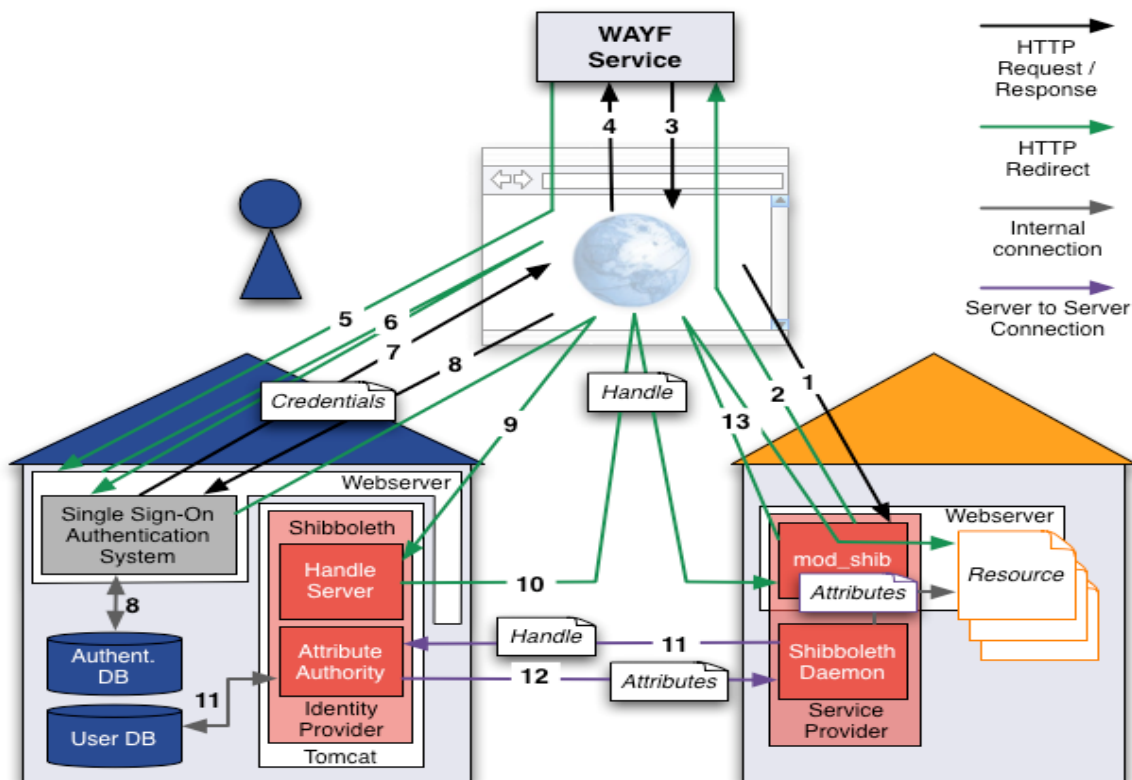


Figure 6: Shibboleth Architecture

In Figure 6 the Shibboleth Architecture is presented. It is composed of two servers and one client. The servers are namely the Service Provider, that is used to secure the resource that the user wants to

access, and the Identity Provider, where the user is authenticated / authorized and the SAML Assertions are created.

Following the above architecture, the TaToo Identity Provider and Directory Service are a Shibboleth Identity Provider installation configured for the TaToo Shibboleth Federation, which is the upper level Shibboleth Federation.

The TaToo Security Frontend is composed of a server that acts as a frontend, that is an Apache HTTPD 2.3 server, where two particular modules, mod_proxy and mod_shib, are running. These modules are respectively responsible for proxying Web Services operations calls toward TaToo Public services and for the Shibboleth business logic concerning the Service Provider. Moreover, the TaToo Security Frontend is composed by a Shibboleth Daemon running on the server host, used to parse SAML Assertions and verify user's attributes.

The last element, the browser, is actually what is defined in TaToo as the User Access Manager. The User Access Manager is a Java library that is used within TaToo portlets or third party clients to authenticate users using the Single Sign-On TaToo infrastructure presented.

In Figure 7 the described Single Sign-On infrastructure is represented with regard to the Shibboleth architecture.

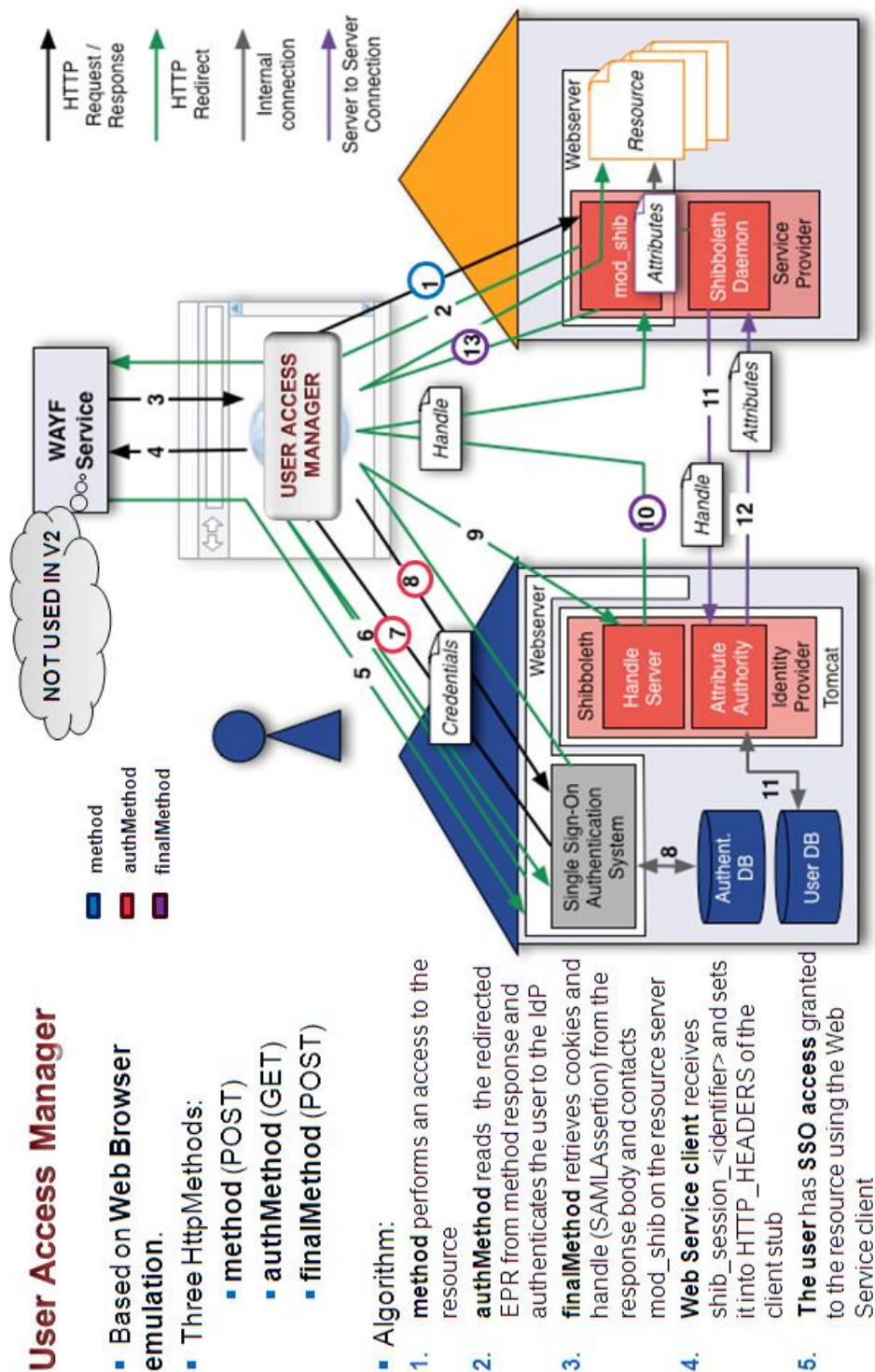


Figure 7: TaToo Single Sign-On infrastructure

2.1.6 Administration

The Administration cross tier as defined in D3.1.3 [2] introduces the concept of manageable components as a means to provide generic administration support in a distributed service network (Figure 8). This concept can be applied to any component, whether it is a User Component, a Public Service or a Core Component. The core concept of this approach is that a manageable component has to provide means to dynamically modify its behaviour, appearance or internal configuration without touching the component implementation itself. Furthermore, a generic approach for the configuration

of manageable components which is independent of the actual configuration methods supported by the component implementations has been established in order to ensure the support for a wide variety of different component types (services, portlets, libraries, etc.).

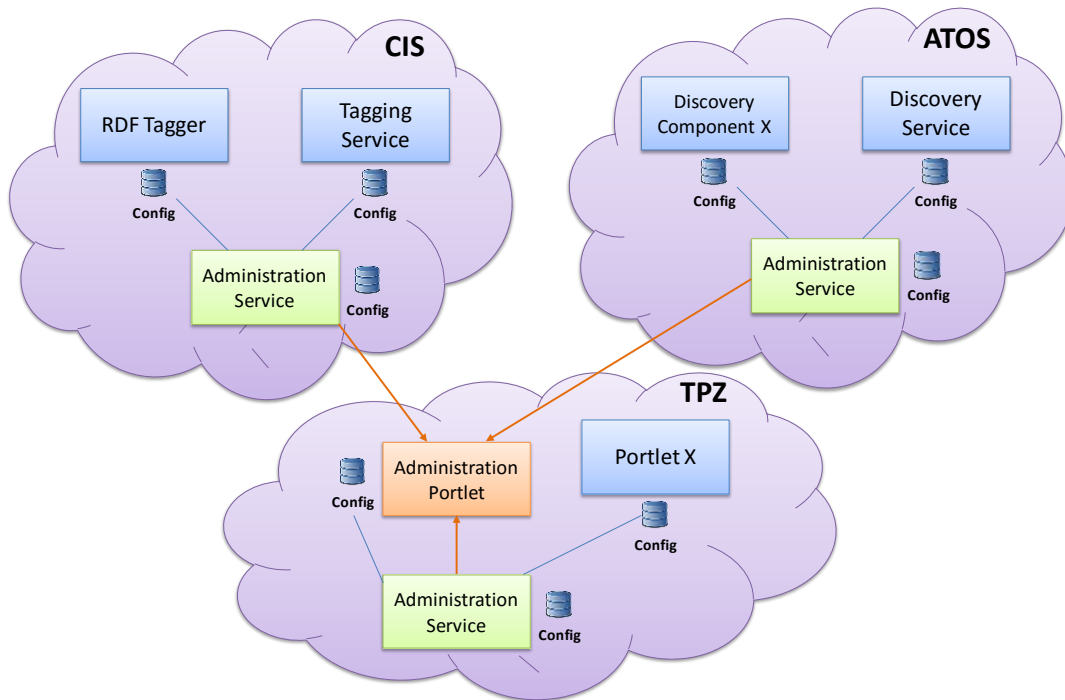


Figure 8: Administration in the TaToo Service Network

The generic and platform-independent approach implemented is based on the Preferences API Specification JSR 10[6], which provides advanced features like change notifications and automatic discovery of new deployed Manageable Components. The Administration Service (see section 2.1.13) shown in Figure 8 above reads and writes the configuration of Manageable Components.

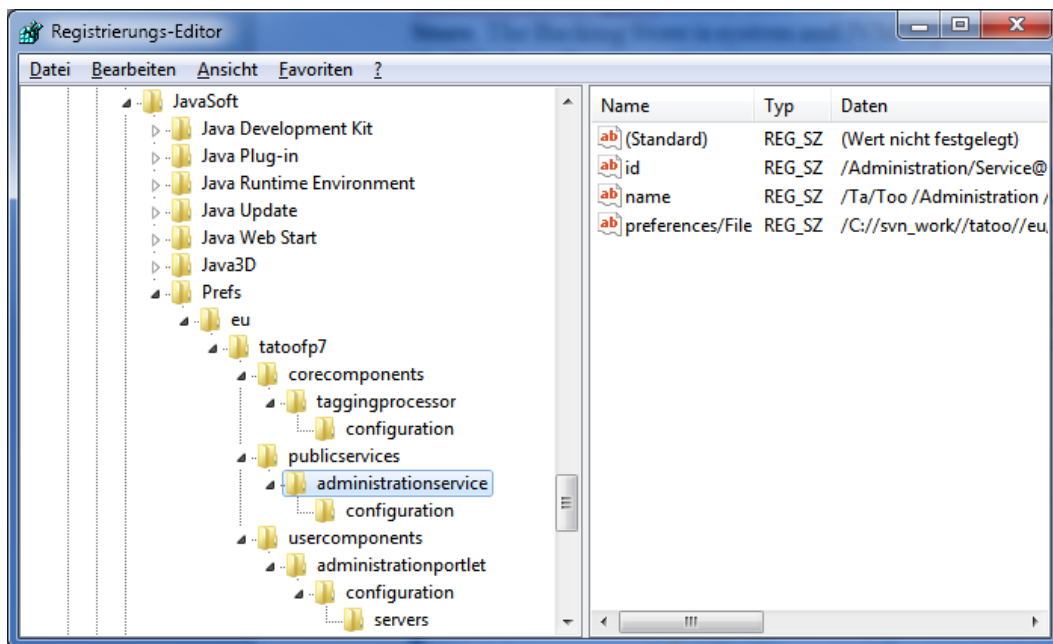


Figure 9: Preferences of the Administration Service (Windows Registry)

The configuration of the component is either stored in a file or, in case of the usage of Java Preferences API, in the so called Backing Store. The Backing Store is an operating system dependent data store for JSR 10 compliant preferences. Under Windows it is implemented by the Windows Registry whereby preferences are stored in HKEY_LOCAL_USER (shown in Figure 9). Under Linux it is represented by a set of xml files stored in the user's home directory (`~/.java/.userPrefs/`).

Access to the Backing Store is transparent to clients of the Preferences API; they do not have to interact with those files directly. To give the administrator the possibility to manually edit the Manageable Components' local configuration without the need to deal with the Backing Store, support for fully automatic export and import of the component's configuration to a local xml file is provided. This logic, as well as the communication with the Java Preferences API is encapsulated in the Administration Support Library, which enables a component developer to integrate instant administration support with minimal effort.

Since the Administration Service has to have local access to the specific Manageable Component's configuration files and the Backing Store respectively, an Administration Service instance has to be present on each server if TaToo Components are deployed on different servers as shown in Figure 8. Thereby, one instance of the Administration Service is able to manage n instances of arbitrary TaToo Components that are enabled for administration support. Furthermore, Manageable Components which make use of the TaToo Administration Support Library can be automatically detected by the Administration Service running on the same server. Administration support of Manageable Components can therefore be gained without any manual configuration. In addition, a component which uses the TaToo Administration Support Library can dynamically react on changes to its configuration thanks to the event notification system of the Java Preferences API. In consequence, TaToo Components can be configured without the need for redeployment or restart.

Graphical editing support of configurations is provided by the Administration Portlet, which is able to communicate with several Administration Service instances. Therefore only one deployed instance of the Administration Portlet is required, to manage all components in a TaToo Service Network. The Administration Portlet is also a Manageable Component by itself which uses the TaToo Administration Support Library, so the list of Administration Services supported by the Administration Portlet can be managed by the Administration Portlet itself.

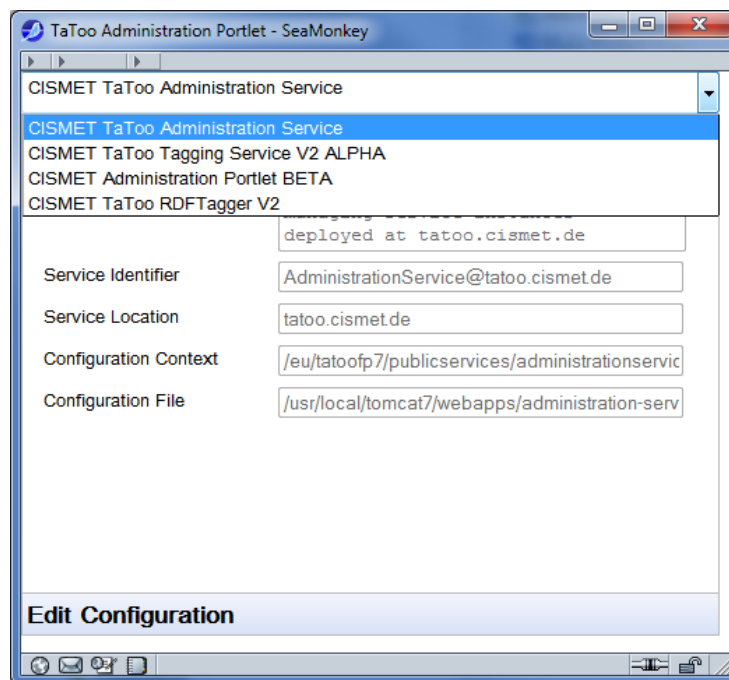


Figure 10: Administration Portlet and supported Manageable Components

2.1.7 Presentation tier components design and implementation description

The presentation tier in TaToo has been provided as Portlets deployed on a Web Portal. The TaToo Web Portal has been implemented adopting the Liferay technology.

Liferay is the world's leading open source enterprise portal solution using latest Java and Web 2.0 technologies. The main characteristics are:

- Runs on all major application servers and web application containers, databases, and operating systems, with over 700 deployment combinations;
- JSR286 compliant;
- Built-in Content Management System (CMS) and Collaboration Suite;
- Personalised pages for all users, both public and private;
- Benchmarked as among the most secure portal platforms using LogicLibrary's Logiscan suite.

Liferay provides also a highly granular permissions system that allows the administrator to customise user experience at organisational and personal level, introducing the concept of Communities. Moreover, Single Sign-On authentication and authorization mechanism is supported.

Within the available options Apache Tomcat application container server has been selected (even if Glassfish is being evaluated), and the portal has been deployed on top of it. The versions used are Apache Tomcat 6.0.26 and Liferay 6.0.5 edition. The portal also needs a production environment database to handle the context automatically, PostgreSQL version 9.0 has been adopted for this purpose.

One of the primary ways of extending the functionality of Liferay Portal is by the use of plug-ins. Portlets are small web applications that run in a portion of a web page. The heart of any portal implementation is its portlets, thus Liferay's core is a portlet container. The container is able to manage portal's pages and aggregate the set of portlets that appear on any particular page. This means that all of the features and functionality of the TaToo Web Portal must reside in its portlets.

The TaToo Portal has been enriched and redesigned in the final version in order to be more user-friendly. A custom TaToo theme and a custom TaToo layout are implemented in order to realise a harmonised experience to the end user accessing the portal. The development of both theme and layout has been realised using development tools provided by Liferay IDE plug-in.

The new approach adopted for the TaToo Portal final version consists in the simplification of the user experience. This is a critical aspect for a portal: if the user is able to take advantage of the TaToo functionality in an easy and comfortable way, the TaToo Portal can attract more and more people to participate and collaborate. It must be noted that the huge participation of users increasing the amount of the TaToo Knowledge Base is one of the main goals for the project.

In particular the new TaToo Portal layout removes the diversification of Tagging, Discovery, and Evaluation functionalities in different portal web pages, and harmonise them in a unique page. This is made possible by an extensive usage of inter-portlet communications called Portlet Events.

In Figure 11 the new TaToo Tools web page is presented. A new intercommunication between portlets is realised so that discovered resources can be tagged "on the fly" by the user, taking advantage of the Tagging Portlet. Discovered Resources can be also evaluated, together with their annotations through the Evaluation portlets implemented during the second cycle.

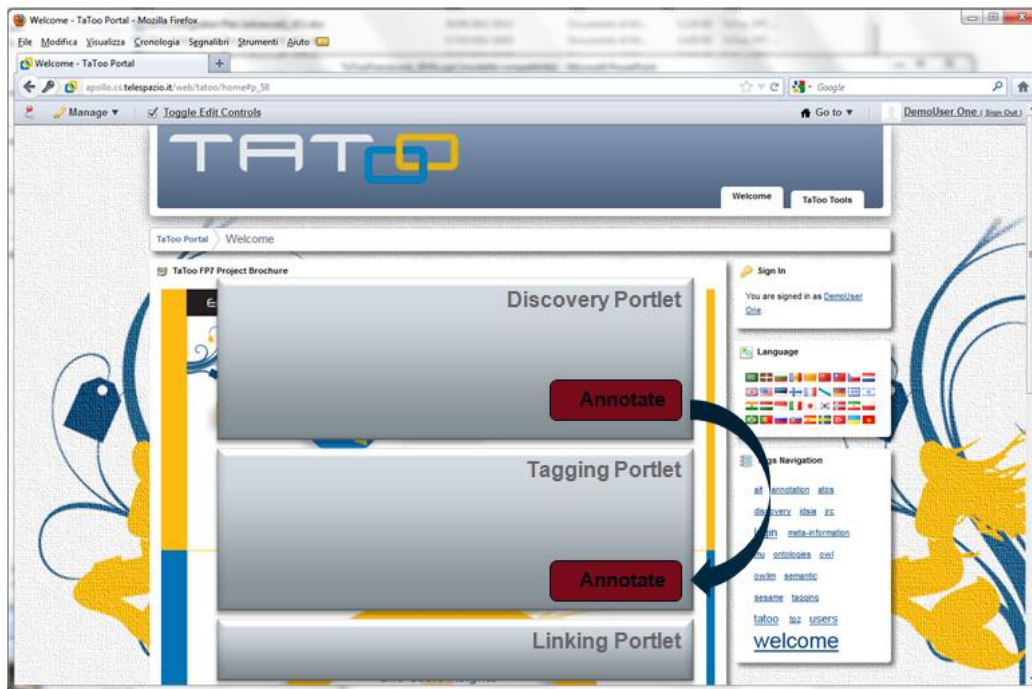


Figure 11: TaToo Portal, new TaToo unique web page

In the final implementation phase the TaToo Portal has been also enriched by the Linking Portlets. The approach that has been adopted is to render it as a pop-up portlet, so that the user is guided while providing a linking, see Figure 12 below.



Figure 12: Evaluation Portlet pop-up.

In this way, the user provides the evaluation of a resource or an annotation triggering the Evaluation portlet from the Discovery Portlets.

The presented approach in the TaToo Portal has been possible also thanks to the adoption of a new and more convenient way to implement TaToo portlets. TaToo portlets are now taking advantage of the

Vaadin framework. Vaadin is a Java framework for building modern web applications, with an appealing look and feel and user friendly user interface. In contrast to pure Javascript libraries and browser plug-in based solutions, it provides a robust server side architecture. The largest part of the application logic thus is actually running on the server instead that in the user browser.

Vaadin framework client (browser) side is based over the Google Web Toolkit (GWT)[7] to enrich the user experience by providing a collection of User Interface components. With respect to pure GWT approach portlet development, Vaadin respects the JSR standards 168 and 286, moreover it is fully integrated with Liferay technology, and provides a good number of useful development and integration tools.

The usage of Vaadin enhances the harmonisation of the specific look and feel of TaToo portlets. Indeed a common Cascading Style Sheets (CSS) file has been designed and adopted by all portlets. In this way the look and feel of the User Interface (portlets) components is harmonised, as well as the fonts, the size and the colours. The adoption of a common CSS file, in conjunction with Liferay and Vaadin layout customisation tools, brings in a new technique that natively supports also mobile browser, providing a new experience to the end user.

TaToo Tagging, Search and Discovery, Evaluation, and Linking Portlets can be deployed in two alternative ways on the TaToo Web Portal. The options provided by Liferay are:

- Local: deploy the portlet directly in the application container locally on the host machine;
- Remote: take advantage of the Liferay control panel which can handle plug-ins installation.

In the context of the TaToo project, being the portal installed on a single site from the first implementation phase, it is foreseen that the remote approach to allow WP4 (Implementation) partners and Use Case partners to deploy their own portlets would be the more convenient one.

Accessing the Liferay control panel with the proper rights, the user can see what plug-ins are currently installed, as shown in Figure 13:

The screenshot displays the Liferay Plugin Installer interface. At the top, there's a 'Plugin Installer' header with tabs for 'Browse Repository', 'Upload File', 'Download File', 'Configuration', and a 'Back' button. Below this, there are tabs for 'Portlet Plugins', 'Theme Plugins', 'Layout Template Plugins', 'Hook Plugins', and 'Web Plugins'. A search section includes a 'Keywords' input field, 'Tag' and 'Repository' dropdown menus (both set to 'All'), an 'Install Status' dropdown (set to 'Not Installed or Out of Date'), and a 'Search Plugins' button. Below the search section, it shows 'Showing 1 - 20 of 84 results.' and pagination controls. A table lists the installed plugins:

Portlet Plugin	Trusted	Tags	Installed Version	Available Version	Modified Date
WOL 5.2.2.1 ID: liferay/world-of-liferay/5.2.2.1/war This is the World of Liferay portlet that integrates social networking features.	Yes	wol	-	5.2.2.1	2/27/09 5:49 PM
Web Form 5.2.2.1 ID: liferay/web-form/5.2.2.1/war	Yes	web	-	5.2.2.1	2/27/09 5:49 PM

Figure 13: Liferay plug-in installation

The user can then install its own developed portlet, packaged as a WAR file, from the upload file section. Once the portlet has been installed on the portal, the user can access their public or private page and add the portlet to it, so that the portlet is accessible and the functionality is available. A summary of the existing TaToo user components and tools is the following:

- Tagging tools:
 - Tagging portlet; It allows the user to browse the topics of an ontology pertaining to its domain and select them to annotate one or more resources identified by an URI.
 - Tags editing portlet: This portlet allows the user to manipulate its previously provided annotations (update and delete annotations).
 - Simple tagging portlet: The portlet is a revised version of the Tagging Portlet, in particular Semantic structures like triple statement composure are hidden to the user. In this way a user can take advantage of TaToo Tagging with no knowledge of Semantics.
 - Geotagging portlet: The Geotagging portlet provides functionality to graphically add location-based annotations supported by the NUTS and GeoNames ontologies to the resources. Also, it shows the location (both point and polygon locations) of discovered resources.
- Discovery tools:
 - Hierarchical search portlet; It allows entering the search criteria using a category tree.
 - Simple search portlet; This portlet provides a form-based functionality for entering search criteria by allowing the user entering controlled metadata about resources and annotations.
 - Results Presentation portlet; The Results Presentation portlet shows the list of resources, the annotations of a selected resource and the RDF triples associated to a selected annotation, and has connections to the tagging, evaluation and linking functionalities.
 - SPARQL portlet; This portlet gives the possibility to query directly the TaToo Knowledge Base. This functionality is provided for advanced users and clients.
- Other tools:
 - Evaluation search portlet; This portlet gives the user the possibility to evaluate annotations and resources.
 - User Access Manager; This client functionality provides an API to seamlessly manage the authentication and authorization of the TaToo portal and public services alike.
 - Administration portlet; It provides access to certain configuration properties of the TaToo Portal.
 - Linking tools; This portlet provide the possibility of linking TaToo resources to external resources in a Linked Data fashion.

2.1.8 Services tier design and implementation description

Public Services in TaToo are mainly implemented as SOAP services. Public services are the entry point of developers to the TaToo functionality. For this reason, this section provides a summary of the services and the methods signature in order to facilitate the developers to get a quick overview of the TaToo API.

2.1.9 Tagging Services design and implementation description

The Tagging Service exposes the tagging functionality of the TaToo System through its public interfaces.

The Tagging Service is implemented as a SOAP web service that allows external clients, in particular the Tagging Portlets, to access the tagging functionality offered by the TaToo System and exposed through the public interface of this service. The current implementation of the Tagging Service receives tagging requests from the different User Components to:

- register new resources
- associate annotations with resources
- retrieve annotations associated with resources
- delete annotations

In V1 the functionality of the Tagging Service was limited to the creation of new tags and the read-only access to existing tags. Furthermore tagging was limited to simple semantic tagging which means choosing from a number of terms of the selected ontology. More sophisticated tagging possibilities, the delete operations and support for more complex MERM compliant Annotations were added in V2.

Major updates on the definition of tagging related data types (resource, annotation, etc.) to support tagging with more complex annotation types were performed in V3.

The Tagging Service interface includes all operations related to tagging. It defines the following operations in Table 2:

Table 1: Tagging service operations

Operation Name	Description
<i>addAnnotationsToResources</i>	This operation is used to add n annotations to exactly n resources. <i>String addAnnotationsToResources(List<BasicResource> resources, List<BasicAnnotation> annotations, String locale) throws MissingParameterValue, InvalidParameterValue, TaTooInternalError, MalformedAnnotation</i>
<i>addAnnotationToResource</i>	This convenience operation is used to add one annotation to one resource. <i>String addAnnotationToResource(BasicResource resource, BasicAnnotation annotation, String locale) throws MissingParameterValue, InvalidParameterValue, TaTooInternalError, MalformedAnnotation</i>
<i>getAnnotationsOfResource</i>	Retrieves all annotations of one resource. <i>List<BasicAnnotation> getAnnotationsOfResource(String resourceURI, String filterOptions, String locale) throws MissingParameterValue, InvalidParameterValue, TaTooInternalError, ResourcesNotFound</i>
<i>getAnnotationsOfResources</i>	Retrieves all annotations of a list of resources. Annotations may be filtered by an optional filter condition, e.g. by a user. <i>Map<String, BasicAnnotation[]> getAnnotationsOfResources(List<String> resourceURIs, String filterOptions, String locale) throws MissingParameterValue, InvalidParameterValue, TaTooInternalError, ResourcesNotFound</i>
<i>removeAnnotation</i>	This operation is used to remove a particular annotation. <i>String removeAnnotation(BasicAnnotation annotation, String foafOnlineAccount) throws eu.tatoofp7.commons.exceptions.MissingParameterValue, eu.tatoofp7.commons.exceptions.InvalidParameterValue, eu.tatoofp7.commons.exceptions.TaTooInternalError, eu.tatoofp7.commons.exceptions.NotAuthorised</i>
<i>removeAnnotations</i>	This operation is used to remove several annotations. <i>String removeAnnotation(List<BasicAnnotation> annotations, String foafOnlineAccount) throws eu.tatoofp7.commons.exceptions.MissingParameterValue, eu.tatoofp7.commons.exceptions.InvalidParameterValue, eu.tatoofp7.commons.exceptions.TaTooInternalError, eu.tatoofp7.commons.exceptions.NotAuthorised</i>

2.1.10 Discovery Services design and implementation description

The Discovery Service exposes the public discovery functionality to User Components. It supports (semantic) search and discovery of annotated resources. The discovery process is query driven, allowing the user to selected certain terms from an ontology that are then used for the semantic search. The Discovery Service itself does not implement any business logic.

The Discovery Service interacts with the Clearinghouse and the search User Component (the Search Portlet). It receives search requests in a specific format from the User Component and transforms them into a message suitable for the Clearinghouse. It receives search results in a specific format from the Clearinghouse and transforms them into a format suitable for the User Component before sending them back.

It is worth noticing that the discovery service provides two kinds of distinct operations:

- **Operations that perform a complete discovery workflow** retrieving as a result a set of complex resource objects (including the resource annotations) that match the original query. The operations in v3 of the discovery service supporting the entire discovery workflow are:
 - **search**: This is the main discovery operation, allowing the client to filter by topic, by annotation types, and by several resource and annotation metadata according to the TaToo ontology framework. The client may also decide whether the search will also retrieve similar resources and/or decide if they would like to perform cross-domain search.
 - **rectangleSearch**: Main operation for performing GeoNames location-based queries within a rectangle delimited by two GeoNames coordinates. This operation is extended by the operation **searchInRectangle**, which combines the previous functionality with extra filtering parameters (the same as in the *search* operation).
 - **nutsSearch**: As in the previous case, this operation is intended to perform location-based queries, but based on NUTS regions (from the NUTS ontology). Again, the **searchWithinNUTS** operation extends *nutsSearch* with extra filtering parameters (the same as in the *search* operation).
- **Operations that give access to calls** (via the Clearinghouse service) **to methods of the three discovery core components, allowing clients to perform a partial execution of the discovery** and retrieving different outputs (resource URIs, annotation URIs, ranked resources, etc.).

The rest of the operations not mentioned above are all of this type. By using in combination several of these methods, clients of the discovery service will be able to tailor their own discovery workflow. As a matter of example, a client might retrieve a set of non-ranked annotations URIs by using the *getGeoAnnotationsInRectangle* method. After doing some processing the client could use the *getResourceOfAnnotations* method to retrieve a non-ranked set of resource URIs that matches with those annotations. With those URIs, the client might perform some more processing or proceed for instance to retrieve a ranked set of resources using one of the ranking methods. After that the client might decide to retrieve the complete resource object and annotation metadata for one or several resources using the *retrieveResources* method.

The Discovery Service is a public interface implemented as a SOAP service that allows external clients to access to the main discovery functionality. It defines the following operations in Table 3:

Table 2: Discovery service operations

Operation Name	Description
search	<i>Main search method implementing the entire discovery workflow. The user might specify if the search should also perform cross-domain queries and expansion of the results by using similarity-based search using the Boolean parameters provided to that effect. The search retrieves a list of complete Resource objects (resource metadata and the set of annotation and their metadata)</i>
	<i>List<Resource> search (boolean searchCrossDomain, boolean extend, String publisherURI, String fromPublishDate, String toPublishDate, String resProvID, String fromResProvDate, String toResProvDate, String annoProvID, String fromAnnoProvDate, String toAnnoProvDate, List<String> topics, String annotationType, String lang, int pageNumber, int pageAmount)</i>
rectangleSearch	<i>Geo-search based in a rectangle query retrieving the complete Resource objects</i>
	<i>List<Resource> search (float lat1, float long1, float lat2, float long2, String lang, int pageNumber, int pageAmount)</i>

Operation Name	Description
searchInRectangle	Simple search and geo-search based in a rectangle query combination (to filter both for location and other metadata) <i>List<Resource> search (float lat1, float long1, float lat2, float long2, String publisherURI, String fromPublishDate, String toPublishDate, String resProvID, String fromResProvDate, String toResProvDate, String annoProvID, String fromAnnoProvDate, String toAnnoProvDate, List<String> topics, String annotationType, String lang, int pageNumber, int pageAmount)</i>
nutsSearch	Geo-search based in a NUTS region retrieving the complete Resource objects <i>List<Resource> search (String nutsURI, String lang, int pageNumber, int pageAmount)</i>
searchWithinNUTS	Simple search and geo-search based in a NUTS region query combination (to filter both for location and other metadata) <i>List<Resource> search (String nutsURI, String publisherURI, String fromPublishDate, String toPublishDate, String resProvID, String fromResProvDate, String toResProvDate, String annoProvID, String fromAnnoProvDate, String toAnnoProvDate, List<String> topics, String annotationType, String lang, int pageNumber, int pageAmount)</i>
simpleSearch	SimpleSearch invoking only the Query expansion functionality (retrieving just a set of resource URIs) The method receives a set of search criteria (resource annotation metadata and topics) and retrieves a list of Resource URIs matching directly the query <i>List<String> simpleSearch (String publisherURI, String fromPublishDate, String toPublishDate, String resProvID, String fromResProvDate, String toResProvDate, String annoProvID, String fromAnnoProvDate, String toAnnoProvDate, List<String> topics, String annotationType, String lang)</i>
crossDomainSearch	Simple search with the extension to cross domain search, invoking only the Query expansion functionality (retrieving just a set of resource URIs) <i>List<String> simpleSearch (String publisherURI, String fromPublishDate, String toPublishDate, String resProvID, String fromResProvDate, String toResProvDate, String annoProvID, String fromAnnoProvDate, String toAnnoProvDate, List<String> topics, String annotationType)</i>
getGeoAnnotationsInRectangle	Invoking directly the equivalent Query expansion functionality. Retrieves a list of resource URIs that have annotations in a rectangle giving two N-S latitude-longitude points of a diagonal of the rectangle <i>List<String> getGeoAnnotationsInRectangle (float lat1, float long1, float lat2, float long2)</i>
getGeoAnnotationsInNUTSRegion	Invoking directly the equivalent Query expansion functionality. Retrieves a list of resource URIs that have annotations in a given NUTS region <i>List<String> getGeoAnnotationsInNUTSRegion(String nutsURI)</i>
getResourceOfAnnotations	For a list of annotations, retrieves the list of resource URIs to what those annotations belongs to. This is useful when the client knows a list of annotation such as the ones provided by the getGeoAnnotation... search methods. Invoking directly the equivalent Query Expansion functionality <i>List<String> getResourceOfAnnotations (List<String> annotationURI)</i>
retrieveResources	Invoking directly the equivalent Results Retriever functionality. Useful when a user knows a set of non-ranked resource URIs list

Operation Name	Description
	<i>and wants to retrieve their complete resource object not ranked.</i>
<i>List<Resource> retrieveResources (List<String> resources, String lang)</i>	
retrieveRankedResources	<i>Invoking directly the equivalent Results Retriever functionality. Useful when a user knows a set of resource URIs and wants to retrieve a ranked by evaluation set of complete resource objects.</i>
<i>List<Resource> retrieveRankedResources (List<RankedResource> resources, String lang)</i>	
rankSearch	<i>Invoking directly the equivalent Results Expansion functionality. Useful to rank either by topic or annotation type a set of known resources (for instance after a simpleSearch is invoked, a subset of the resources can be sent to this method for ranking)</i>
<i>List<RankedResource> rankSearch (List<String> resultResources, List<String> topics, String annotationType)</i>	
rankByEvaluation	<i>Invoking directly the equivalent Results Expansion functionality Useful to rank by evaluation a set of resources and their evaluations.</i>
<i>List<RankedResource> rankByEvaluation (List<String> resultResources, Evaluation[][] resultEvaluations)</i>	
expandResultsThroughSimilarResources	<i>Invoking directly the equivalent Results Expansion functionality. Useful to expand the results with similar resources (similarity-based search) once the user has a list of valid resources.</i>
<i>List<String> expandResultsThroughsimilarResources(List<String> resources)</i>	
getAnnotationsByCategory	<i>Retrieves URIs of annotation instances which forms a triple with the propertyURI and categoryURI respectively</i>
<i>List<String> getAnnotationsByCategory(String propertyURI, String categoryURI)</i>	

2.1.11 Evaluation Services design and implementation description

The Evaluation Service is designed as a SOAP Web Service. The goal of this service is to enable external clients (in particular the Evaluation Portlet) to access the evaluation functionalities provided by the TaToo Framework. The Evaluation Processor implements those functionalities. The Evaluation Service is specified according to the service implementation specification template and the TaToo SOAP Web Service Platform.

The implementation of the service is done using JAX-WS API and the service is deployed on Apache Tomcat 7 application container. The Evaluation Service accesses the TaToo evaluation functionalities through the Clearinghouse service. It defines the following operations in Table 4:

Table 3: Evaluation service operations

Operation Name	Description
addResourceEvaluation	This operation adds a resource evaluation determined by evaluation criterion, metric, value, creation date, evaluator, short description, and a description language. It is also responsibility of this operation to update a resource's evaluation score that corresponds to the evaluation criterion of the added evaluation.
<i>String addResourceEvaluation (String resourceURI, Evaluation evaluation)</i>	
addAnnotationEvaluation	This operation adds an annotation evaluation determined by evaluation criterion, metric, value, creation date, evaluator, short description, and a

Operation Name	Description
	description language. It is also responsibility of this operation to update an annotation's evaluation score that corresponds to the evaluation criterion of the added evaluation.
<i>String addAnnotationEvaluation (String annotationURI, Evaluation evaluation)</i>	
getEvaluationOfResoure	This operation retrieves all existing evaluations of a resource specified by the resource's URI.
<i>Evaluation getEvaluationOfResource (String resourceURI, String filterOptions)</i>	
getEvaluationOfAnnotation	This operation retrieves all existing evaluations of an annotation specified by the resource's URI.
<i>Evaluation getEvaluationOfAnnotation (String annotationURI, String filterOptions)</i>	
getEvaluationsOfResources	This operation retrieves all existing evaluations of a list of resources specified by their URIs.
<i>HashMap<String, List<Evaluation>> getEvaluationsOfResources (List<String> resourceURIs, String filterOptions)</i>	
getEvaluationsOfAnnotations	This operation retrieves all existing evaluations of a list of annotations specified by their URIs.
<i>HashMap<String, List<Evaluation>> getEvaluationsOfAnnotations (List<String> annotationURIs, String filterOptions)</i>	
removeResourceEvaluation	This operation removes a resource evaluation specified by the evaluation's URI. It is also responsibility of this operation to update the resource's evaluation score that corresponds to the evaluation criterion of the removed evaluation.
<i>String removeResourceEvaluation (String evaluationURI, String userToken)</i>	
removeEvaluationsOfResoure	This operation removes all evaluations of a resource specified by the resource's URI. It also removes all the resource's evaluation scores.
<i>String removeEvaluationsOfResource (String resourceURI, String userToken)</i>	
removeAnnotationEvaluation	This operation removes an annotation evaluation specified by the evaluation's URI. It is also responsibility of this operation to update the annotation's evaluation score that corresponds to the evaluation criterion of the removed evaluation.
<i>String removeAnnotationEvaluation (String evaluationURI, String userToken)</i>	
removeEvaluationsOfAnnotation	This operation removes all evaluations of an annotation specified by the annotation's URI. It also removes all the annotation's evaluation scores.
<i>String removeEvaluationsOfAnnotation (String annotationURI, String userToken)</i>	
getResourceEvaluationScore	This operation retrieves the resource's evaluation score for the given evaluation criterion.
<i>Double getResourceEvaluationScore (String resourceURI, String evaluationCriterion)</i>	
getAnnotationEvaluationScore	This operation retrieves the annotation's evaluation score for the given evaluation criterion.
<i>Double getAnnotationEvaluationScore (String annotationURI, String evaluationCriterion)</i>	

2.1.12 Ontology Manager Services design and implementation description

The Ontology Manager Service is a public interface implemented as a SOAP service that allows external clients to access to the TaToo semantic repositories to retrieve ontology elements. The Ontology Manager offers functionality to retrieve filtered information about ontologies from the

Semantic Repository. Therefore it provides supporting functionality for clients of the TaToo Framework.

The Ontology Manager interface defines the following operations Table 5:

Table 4: Ontology Manager service operations

Operation Name	Description
getOntology	This operation retrieves an ontology framework for the specified domain given as input.
<i>String getOntology(URI domain)</i>	
listDomains	This operation retrieves the list of domains via their assigned URIs (contexts)
<i>List<URI> listDomains()</i>	
getPrefixNS	This operation retrieves the list of available namespaces of the ontologies in the TaToo knowledgebase
<i>List<String> getPrefixNS()</i>	
getMERMOntology	This operation retrieves the MERM ontology in RDF/XML format
<i>String getMERMOntology()</i>	
sparqlQuery	This operation executes a SPARQL given as input and returns the result in SPARQL query results XML format ¹³ .
<i>String sparqlQuery(String query)</i>	
getResourceProviders	This operation retrieves all the provider accounts that at least provided one resource.
<i>Set<Account> getResourceProviders()</i>	
getAnnotationProviders	This operation retrieves all the provider accounts that at least provided one annotation
<i>Set<Account> getAnnotationProviders()</i>	
getResourcePublishers	This operation retrieves all the provider accounts that at least provided one annotation
<i>Set<Account> getAnnotationProviders()</i>	
getAnnotationTypes	This operation retrieves all types of annotation in the specified domain
<i>Set<Type> getAnnotationTypes(URI domain, String lang)</i>	
getSubjectTypes	This operation retrieves all types of links semantically available between the specified annotation type and topic in the specified domain
<i>SetOfURI getSubjectTypes(URI domain, URI topic, URI annotationType)</i>	
getTopics	This operation retrieves all types of topics in the specified domain
<i>Set<Topic> getTopics(URI domain, String lang)</i>	

2.1.13 Administration Service design and implementation description

The Administration Service is a Public Service implemented as SOAP web service that manages Manageable Components. It does not interact with the components to be managed directly. It either accesses a local configuration file of the component or makes use of the Java Preferences API (JSR 10) to update the configuration of a component. The advantage of the usage of the Java Preferences API is that a component can register an event listener to dynamically react to configuration changes made by the Administration Service. The Administration Service in turn is able to automatically detect new components by watching the system Backing Store for new component configurations.

The Administration Service defines the following operations in Table 6:

¹³ <http://www.w3.org/TR/rdf-sparql-XMLres/>

Table 5: Administration Service operations

Operation Name	Description
<i>listManageableComponents</i>	Returns the manageable components supported by this service instance.
<i>ManageableComponent[] listManageableComponents() throws TaTooInternalError</i>	
<i>getConfiguration</i>	Retrieves the configuration of a manageable component identified by its id.
<i>String getConfiguration(String componentId) throws MissingParameterValue, InvalidParameterValue, TaTooInternalError</i>	
<i>updateConfiguration</i>	Updates the configuration of a manageable component.
<i>void updateConfiguration(String componentId, String configuration) throws, MissingParameterValue, InvalidParameterValue, TaTooInternalError</i>	

2.1.14 Linking Service design and implementation description

The Linking Service is designed as a SOAP Web Service. The goal of this service is to enable external clients (in particular the Linking Portlet) to add a set of typed links between the TaToo resources as well as to interlink similar resources by adding similarity links between them. The Linking Service is specified according to the service implementation specification template and the TaToo SOAP Web Service Platform.

The implementation of the service is done using JAX-WS API and the service is deployed on Apache Tomcat 7 application container. The Evaluation Service accesses the TaToo linking functionalities through the Clearinghouse service. It defines the following operations:

Table 6: Linking service operations

Operation Name	Description
<i>addLinks</i>	This operation creates links between a given resource and a list of resources related to it. The links' type is determined by a given link property.
<i>String addLinks (String resourceURI, String Linkproperty, List<String> ListOfresourceURIs)</i>	
<i>addSimilarityLinks</i>	This operation creates similarity links between a given resource and a list of resources similar to it. The links' type is determined by the TaToo resource similarity property defined in MERM.
<i>String addSimilarityLinks (String resourceURI, List<String> ListOfresourceURIs)</i>	
<i>getLinkedresources</i>	This operation retrieves all resources linked to a given resource regardless of the link types.
<i>List<String> getLinkedResources (String resourceURI)</i>	
<i>getLinkedResourcesByGivenRelationship</i>	Mandatory operation that retrieves all resources linked to a given resource by links of a given link type.
<i>List<String> getLinkedResources (String resourceURI, String linkType)</i>	
<i>getSimilarResources</i>	This operation retrieves all resources similar to a given resource.
<i>List<String> getSimilarResources (String resourceURI)</i>	
<i>removeAllLinks</i>	This operation removes all links of the given resource.
<i>String removeAllLinks (String resourceURI)</i>	
<i>removeAllLinksByGivenRelationship</i>	This operation removes all links of the given resource and the given link type.
<i>String removeAllLinksByGivenRelationship (String resourceURI, String linkType)</i>	

<i>removeLink</i>	This operation removes the specified link (i.e., resource1, link type, resource 2).
<i>String removeLink (String resourceURI1, String linkType, String resourceURI2)</i>	

2.1.15 Business tier design and implementation specification

Core Components are the implementation of the TaToo business logic. These components should be carefully implemented considering available technology of success. In particular, Core Components relying on semantics are, in general, implemented by widely adopted existent software components. Those components may require some infrastructure e.g. Sesame must be deployed on a web application server and may use a RDBMS as persistence storage. These infrastructure elements are described next in this section.

2.1.16 Clearinghouse design and implementation specification

The Clearinghouse is a TaToo Core Component that represents the central entry point to the TaToo Business Tier from the TaToo Public Services. All requests that come from the TaToo Public Services and which are addressed to any of the TaToo Core Components must go through the Clearinghouse component.

For the requests that are served in one step by one single core component (e.g., get user profile, get ontology, and get ontology URI), the role of the Clearinghouse is just to delegate the request to the appropriate core component and then send the results back to the public service that initiated the request.

For the requests that are served in several steps (e.g., search and tagging), each of which is realized by a different core component, the Clearinghouse will also provide appropriate workflow logic. For example, in case of the search request, the process workflow is composed of three steps and involves three core components (i.e., the Query Expansion, the Resource Retrieval, and the Result Expansion components). After receiving the search request, the Clearinghouse first employs the Query Expansion component to expand the initial user query, then it employs the Resource Retrieval component, which executes the expanded query against the TaToo Semantic Repository obtaining a set of resources relevant to the query, and at the end it employs the Result Expansion component to discover additional resources related to those that are already retrieved. Besides serving requests that come from the TaToo Public Services, the Clearinghouse also controls the process of harvesting metadata from the external resource. It provides the logic that manipulates the list of external sources holding resources of interest for the TaToo and triggers the TaToo Resource Harvester component.

The Clearinghouse component is implemented as a SOAP web service using JAX-WS API. The design of the internal Clearinghouse logic conforms to well known multithreading and pipeline software design principles[8].

2.2 Semantic Processor design and implementation specification

From the architectural point of view, the Semantic Processor is a Building Block made of a set of Core Components (from the implementation point of view it can be implemented as a single component). It offers the infrastructure to access to the TaToo Semantic Repository and ontologies allowing semantically enhanced tagging and discovery functionality. It supports the tagging and discovery processes by providing functions to:

- retrieve ontologies on the basis of user-defined properties like domain or context;
- store and retrieve semantic annotations for resources (RDF-Triples);
- search for semantic annotations;
- semantic reasoning based on the existing ontologies.

The Semantic Processor provides its functionality to Core Components either directly or through the Clearinghouse, and to Public Services (and User Components in turn) through the Clearinghouse only. The semantic functionality is realised by the following components:

- The **Ontology Manager** to manage access to the TaToo ontologies;
- The **Semantic Framework** to manage the interaction with the Semantic Repository and the Reasoner;
- A **Reasoner** to infer new knowledge from the available RDF-Triples & ontologies and to check for incongruent (or inconsistent) information while managing ontologies;
- A **Triplestore** (Semantic Repository) which is a data base for the storage and retrieval of RDF-Metadata.

2.2.1 Ontology Manager

The Ontology Manager Service is a core component implemented as a SOAP service that allows the Clearinghouse to access to the TaToo semantic repositories to retrieve ontology elements. The Ontology Manager offers functionality to retrieve filtered information about ontologies from the Semantic Repository. The Ontology Manager allows retrieving ontologies and ontology elements in the TaToo Framework. The Ontology Manager has been realised as a component accessible via the Ontology Manager Public Service, as explained in section 2.1.12.

2.2.2 Semantic Framework

The main functionalities of the Semantic Framework, Reasoner and Triplestore (Semantic Repository) components are covered by an existing **Semantic Framework** to take advantage of a set of useful APIs to manipulate RDF, support SPARQL queries and reasoning and RDF storage in a Semantic Repository.

In this sense we have evaluated the usage of two frameworks: Sesame (SESAME) and Jena (JENA). Both frameworks have similar characteristics (APIs, RDF support, etc.), although according to our findings Jena is better in terms of reasoning and Sesame has a better user interface and Semantic Repository capabilities. Besides, the Sesame framework can be extended with OWLIM, a native RDF engine, implemented in Java, which has a good support for the semantics of OWL 2 RL and a proven scalability, and moreover offers out-of-the-box features for issuing geo-queries. Therefore, we have decided to implement the Semantic Framework using Sesame and OWLIM.

The Semantic Framework covers the following functionalities:

- The Semantic Processor functionalities, as the central management of access to, and search for semantic information. It provides functions for semantic search (e.g. by providing a SPARQL endpoint), storage and retrieval of RDF-Triples (e.g. tags), inference through the Reasoner, etc.
- The Semantic Reasoner (semantic inference machine) component is used to perform inference in the TaToo Knowledge Base (based on the existing ontologies and meta-information).
- The RDF Semantic Repository is part of the Data Access Component dealing with RDF storage.

2.2.3 Discovery Processor design and implementation description

Discovery Processors are supporting components realising the search process functionalities. Currently, the following discovery processors are implemented:

- The **Query Expansion** component is responsible for transforming the user's information need into a query executable by the Semantic Processor. Different components may perform this transformation in different ways ranging from a direct conversion to processes that include generalization of terms, resource-specific queries, etc. The Query Expansion Component takes the user query from the Clearinghouse as input and converts it into one or a set of SPARQL queries. Then the component connects to the Semantic Framework to execute the query. Finally it delivers a set of resource or annotations as an output. The Query Expansion provide one main method for filtered search (simpleSearch), and two main methods for geo-search (rectangleSearch -for geo-queries based on GeoNames-, and

nutsSearch -for geo-queries based on NUTS). It provides also methods for using cross-domain search in combination with the previous methods.

- The **Resource Retriever** component is responsible to retrieve and aggregate information relevant to a given query about resources available in the system. Different components may perform this process in different ways, recovering only annotations for a given domain, grouping similar annotations, etc. Its main objectives are therefore i) to retrieve meta-information of specified resources and their annotations, creating as output complex resource objects and ii) to retrieve meta-information of specified annotations, creating as output complex annotation objects. In combination with the Results Expansion service it provides a ranked set of complex resource objects.
- The **Results Expansion** component is responsible for the result ranking and enrichment. Different components may perform this process differently, using different ranking strategies, tailoring the results to the user, obtaining information from external sources, etc. Thus, the component provides methods for further expansion of the main query results (especially taking into account similarity-based aspects) and ranking of the resources.

For the ranking aspects, the Resource Expansion Component may take an annotation or a set of unranked results as an input. It queries all the necessary meta-information from the semantic repository of the annotation or ranks the result based on the evaluation of the resources/annotations, or the relevance to the query (the closest match to the topics provided in the query).

For the expansion of the results, the Resource Expansion takes advantage of the cross-mapping between domain ontologies and the similarity stated between TaToo resources according to the TaToo similarity model.

2.2.4 Tagging Processor design and implementation description

Tagging Processors are supporting components realising certain functionalities of the tagging process. Currently, the following tagging processors are implemented:

- Schema Mapping Component prototype
- Visualization and Filtering Component
- RDF Tagger
- Linking processor

The Schema Mapping Component prototype supports the mapping from one xml-schema to another where the mapping rules are described in XSLT, and it is realised as a SOAP Web Service.

The RDF Tagger is responsible for storing annotations as well as resources in the TaToo knowledgebase in a format that is compliant to the MERM ontology. It provides the actual business logic of the Tagging Service and saves annotation and resource objects received from the Tagging Service (through the Clearinghouse) as instances of the MERM ontology whereby it performs validity and plausibility checks of the annotations.

2.2.5 Resource Harvester design and implementation specification

The Resource Harvester (or simply Harvester) is the Core Component responsible for harvesting meta-information from available resources that could be data (catalogues), web services, web pages, etc. A more and more large meta-information set is essential to improve the process of searching for resources returning good results to the requesting user.

For each distinct resource type, the harvesting functionality is realised by a specialised Harvester Connector implemented as a plug-in to the Harvester. A Catalogue Connector, for example, is used to retrieve meta-information stored in catalogues. Since there exist no common meta-information schema for catalogues, a Catalogue Connector has to be provided for each distinct schema that shall be supported. The implementation of a Harvester Connector is the implementation of a component adhering to a Java interface (the ConnectorInterface). Connectors implementing the interface and properly configured are dynamically loaded and used by the Resource Harvester.

As result of the final implementation phase the following Harvester connectors have been implemented and are currently harvesting resource:

- GENESI-DEC Connector - harvests OGC meta-information from GENESI-DEC catalogues, only those were pertaining to TaToo domains are then stored in the repository;
- JRC Connector – harvests meta-information from JRC Validation Scenario catalogues, generated in the context of the TaToo project;
- Masaryk University Connector – harvests meta-information from Masaryk University Validation Scenario, generated from the legacy RDF;
- LinkedData Connector – harvests resources pertaining to the TaToo domains that are part of LinkedData;
- Web Site Connector – harvests meta-information stored as RDFa in any Web Site available. As a proof of concept the TaToo Project Web Site has been harvested and stored as a resource.

In the following sections the connectors developed are described in detail.

2.2.6 Harvester connectors

GENESI -DEC Connector

The GENESI-DEC Connector is a Resource Connector (Catalogue Connector) plug-in for the Resource Harvester component. As foreseen by the Resource Harvester component, it implements the ConnectorInterface Java interface.

In order to harvest the GENESI-DEC catalogues, the component uses the OpenSearch protocol through the Rome (a9 OpenSearch module) / Apache Abdera (OpenSearch extension) libraries. Considering that GENESI-DEC catalogues are able to provide results to OpenSearch queries in RDF format, the connector has no need to convert a different XML based format to RDF (e.g. through GRDDL).

GENESI-DEC Resources meta-information is structured in different levels of catalogues, namely: Datasets, Series, and Datasets Series. Datasets are contained in a unique catalogue containing different Series and Services, where Series are containing a set of Datasets Series and Services are a list of OGC (Open Geospatial Consortium). Finally Dataset Series are a set of Data.

Currently the catalogue of GENESI-DEC exposing Datasets counts more than two thousand resources, however this number grows really fast scaling down to Datasets to a size of million resources. In the context of GENESI Domain the resources that are interesting for discovery and tagging are the Datasets. Series and Datasets Series are subject to authentication / authorization and specific security policies, thus in TaToo the upper level catalogue containing Datasets is harvested. However not all Datasets are stored into the TaToo Knowledge Base repository due to the heterogeneity of data, that goes from Spatial specific data to Marine data. To harvest only Environmental resources that could have cross alignments also with other TaToo Domains, a filtered small set of terms from the GENESI Domain Ontology (composed by existing ontologies, i.e. GCMD or GEMET) has been bridged to the TaToo Bridge Ontology.

The GENESI connector performs a syntactic match on Title, Abstract and Description text areas filled in the catalogue. Thanks to this approach only interesting GENESI resources for TaToo Domains are harvested, currently counting over six hundreds Datasets, that can scale up to thousands Dataset Series.

JRC Harvester Connector

The JRC Harvester Connector is a Resource Connector plug-in for the Resource Harvester component. This connector provides capabilities for harvesting resources provided by the AGRI4CAST unit of the JRC MARS. This unit is centred on the JRC's crop yield forecasting system aiming at providing accurate and timely crop yield forecasts and crop production biomass for the EU territory.

The resources that are provided by AGRI4CAST and which are considered for the TaToo harvesting include: agricultural software libraries and tools, weather data, agricultural maps, remote sensing data, etc.

AGRI4CAST provides the RDF descriptions of the resources to be harvested. As being one of the TaToo Validation Scenarios partner, the AGRI4CAST JRC unit generates the RDF resource descriptions according to the TaToo Semantic Framework (i.e., the MERM and Bridge Ontologies) and applies the JRC ontology for the resource annotation. The JRC ontology is also properly aligned to the TaToo Semantic Framework. Taking all this into account, it is clear that the role of the JRC connector is mainly to access and retrieve the AGRI4CAST resource descriptions from the JRC server, validate them against the TaToo Semantic Framework and store them to the TaToo Knowledge Base (i.e., the TaToo RDF repository).

In order to read the RDF descriptions of the resources as well as to store them to the TaToo RDF repository, the JRC Harvester Connector utilises the Sesame API.

Masaryk University Connector

The Masaryk University Connector is a Resource Connector plug-in for the Resource Harvester component. Resources provided from the Masaryk University are in the context of POPs (Persistent Organic Pollutants) and Oncological Data that are used to study the anthropogenic impact and global climate change influence on the trajectory of Persistent Organic Pollutants.

The Masaryk University Catalogue is exposed using RDF language and is already annotated by concepts taken from the Masaryk University Domain Ontology developed in the scope of the TaToo project.

The MU connector then parses the RDF catalogue using Sesame API and stores directly the RDF triple statements retrieved into the TaToo repository. Created Resources and Annotations are fully compliant with the MERM structure defined in TaToo.

TaToo Web Site Harvester Connector

The TaToo Web Site Harvester Connector is a Resource Connector plug-in for the Resource Harvester component. This connector provides capabilities for harvesting resources provided by the web sites which use RDFa. The harvested web sites embed rich metadata by adding a set of attribute-level extensions to XHTML. This allows RDF data-model mapping and enables the use of RDF subject-predicate-object expressions.

The resources which are considered for the TaToo harvesting are therefore web sites with RDFa elements. In the first phase supported RDFa elements describe the web site as a resource and define the elements “resource name”, “resource description” and “resource uri”.

The Web Site Harvester connector parses the RDFa descriptions from web sites and stores it for further processing in a RDF repository. It generates RDF resource descriptions according to the TaToo semantic framework. Therefore the role of the Web Sites Connector is to access and retrieve resource descriptions from web sites (as well as additional data in RDFa format), validate them against the TaToo Semantic Framework and store them to the TaToo Knowledge Base.

In order to read the RDF descriptions of the resources as well as to store them to the TaToo RDF repository, the Web Site Harvester Connector utilises the Sesame API.

Linked Data Connector

The Linked Data Connector provides harvesting capabilities to the TaToo Framework to crawl and harvest Linked Data resources. The connector is able to integrate TaToo resources annotated in the Knowledge Base repository with linked resources from the Semantic Web. The connector presents then two main advantages, largely increases the knowledge base stored in the TaToo repository, and plugs the TaToo resources in the Linked Open Data cloud.

The Linked Data Connector is based on an already existing framework that crawls the Linked Data Web, namely LDSpider¹⁴. The LDSpider API used in TaToo offers a web crawler adapted to traverse and harvest content from the Linked Data web. The LDSpider is used to perform the first action of the Linked Data harvesting workflow that is crawling resources from the Linked Web. The harvesting process can be modified on demand from the Harvester Service using the following configuration parameters:

- rounds, that is the desired depth of search;
- maximum amount of URIs to be crawled at each round;
- the list of ontologies properties that must be filtered in the result set of triple statements.

The implemented connector filters the result triples based on the link types supported in TaToo (see Linking Processor design and implementation description). Once the Linked Data connector receives the filtered set of triple statements, it directly stores the triples set using the TaToo Linking library.

The final step of the Linked Data workflow is to perform a second round parsing of the harvested LOD (Linked Open Data) cloud datasets to check if there are meaningful annotations in the context of TaToo. In this case a specific mapping is performed and resources are also semantically annotated with concepts from TaToo Domain Ontologies.

The choice of seed has to be carefully taken in consideration, since depending the depth level adopted for the harvesting process, different types of resources will be harvested and stored into the TaToo Frameworks. Currently TaToo is harvesting the Diseasesome Linked Data datasets, where a large number of relationships with the Masaryk University Validation Scenario have been found. In particular, Diseasesome publishes a network of 4,300 disorders and disease genes linked by known disorder-gene associations for exploring all known phenotype and disease gene associations, indicating the common genetic origin of many diseases. This new set of links and relationships enrich the available information in the context of Oncological or POPs resources search. Also, the CIA World Factbook dataset is harvested in relation with the AIT Validation Scenario, enriching information on discovered twin regions from the AIT legacy application search.

2.2.7 Evaluation Processor design and implementation description

The Evaluation Processor is designed as a SOAP Web Service, thus it is specified in accordance to the service implementation specification template and the TaToo SOAP Web Services Service Platform.

The prototype of the Evaluation Service is implemented by using the JAX-WS API and is deployed on the Apache Tomcat 7 application container.

The Evaluation Processor is component of the TaToo business providing necessary functionalities for the evaluation of the TaToo-managed resources and the resource annotations. The functionalities of the Evaluation Processor could be grouped into the following four categories:

- Creating new evaluations of the resources and the resource annotations,
- Retrieving existing resource and annotation evaluations,
- Removing existing resource and annotation evaluations, and
- Generating and retrieving average evaluation scores of the resources and resource annotations.

The Evaluation Processor communicates with the TaToo Knowledge Base to store and retrieve the evaluation information. For that purpose, the Evaluation Processor utilises Sesame API that provides capabilities to read and write RDF triples to the TaToo Knowledge Base as well as to execute SPARQL queries against the TaToo Knowledge Base SPARQL endpoint.

2.2.8 Linking Processor design and implementation description

The Linking Processor is designed as a SOAP Web Service, thus it is specified in accordance to the service implementation specification template and the TaToo SOAP Web Services Service Platform.

¹⁴ <http://code.google.com/p/ldspider/>

The Linking Processor is implemented by using the JAX-WS API and is deployed on the Apache Tomcat 7 application container.

The Linking Processor is a component of the TaToo Business tier, which implements the TaToo linked data functionalities. In particular, it provided the following list of functionalities:

- Adding a predefined set of typed links between TaToo resources;
- Adding similarity links between similar TaToo resources;
- Retrieving a list of resources linked to a given resources;
- Retrieving a list of similar resources of a given resource;
- Removing the typed links and similarity links set between TaToo resources.

The Linking Processor communicates with the TaToo Knowledge Base to store, retrieve, and remove typed and similarity links. For that purpose, the Linking Processor utilises Sesame API that provides capabilities to read and write RDF triples to the TaToo Knowledge Base as well as to execute SPARQL queries against the TaToo Knowledge Base SPARQL endpoint.

2.2.9 Data Access Component design and implementation specification

The reference implementation of the Access Component to the TaToo Semantic Repository uses existing software (Sesame API and SPARQL). See more details in section 2.2.

2.3 Ontology development

The TaToo project aims at semantic tagging and searching environmental resources. TaToo aims to capitalise on the principles of the Semantic Web using ontologies as the underlying model for tagging and searching resources. The basis for the usage of ontologies within the TaToo Framework has been explained in the deliverable D3.1.3[2]. The main objectives of the ontologies within the TaToo Framework are the following:

- Allow formal tagging and searching of environmental resources.
- Allow contextual cross-domain tagging and searching of environmental resources.
- Foster multilingualism issues.
- Implement an extensible and more accurate searching mechanism.

In order to achieve these objectives, the conclusions achieved since the delivery of D3.1.1 and updates in D3.1.3, the premises for the TaToo ontologies are the following:

- Usage of standards: TaToo ontologies are based on existing W3C standards, particularly RDF [9], RDFS[10] and OWL[11], in particular a subset of OWL 2 called OWL2 RL. In the ontology engineering process we have also studied and reused many shared vocabularies, such as Dublin Core[12], FOAF[13] and SIOC[14].
- Usage of methodology: Within TaToo the use of the NeOn methodology[15] has been fostered, as decided in deliverable D3.1.1[1].
- Reuse of existing ontology engineering tools: TaToo will not provide an ontology engineering tool, but instead will rely on existing tools for ontology engineering, such as the NeOn Toolkit[16] or Protégé[17]. The NeOn Toolkit is the preferred choice as it gives tooling support for the selected methodology.
- Avoid the use of complex semantic for tagging and discovery

2.3.1 Conceptualization and formalization of TaToo ontologies

The TaToo Ontology Framework

As explained in D3.1.3[2] the NeOn methodology for engineering ontologies was adopted[18]. This methodology provides guidance for all aspects of the ontology development, focusing on collaborative aspects on designing ontologies, the reuse of existing resources (ontological and non-ontological), and

the further maintenance and evolution of the ontologies. The NeOn Toolkit includes tooling support (plug-ins) supporting some of the activities of the methodology.

Based on the methodological guidelines from the NeOn methodology, the work on ontologies within TaToo has followed several iterations and used multiple methodological scenarios. The TaToo Ontology Framework implementation methodology and steps followed has been amply described in previous deliverables (TaToo D411, 2011). A summary of the main steps and milestones is the following:

- The design and implementation of the TaToo ontologies have followed the NeOn methodology, especially scenarios for ontology reuse and ontology re-engineering;
- Gathering of competency questions in several workshops with Validation Scenario users. The main conclusions drove towards the definition and implementation of a common minimal model to describe environmental resources and annotations, and a semantic infrastructure where different domain ontologies could be relatively easy plugged in;
- Need of a model for evaluation of the resources.

After the collection of requirements, three main issues impacted on the shape of the TaToo Ontology Framework, namely: i) take into account **multilingual aspects**; ii) allow **formal tagging and searching** environmental resources; and iii) allow **cross-domain search**.

With respect to the **multilingual approach** of the TaToo ontologies, Deliverable D3.1.3[2] provides an ample rational. The final decision made in TaToo is to use the RDFSchema metadata `rdfs:label` to add multilingual metadata storage in the TaToo Ontologies (including Domain Ontologies).

Formal tagging and searching of heterogeneous and disperse environmental resources implies to define an ontology resource model. For practical reasons formal tagging in opposition to informal tagging should use a common and well-defined minimal model. For this purpose the Minimum Environmental Resource Model (MERM) was introduced in deliverable D3.1.1[1]. MERM is defined as the largest common denominator between a set of heterogeneous description formalisms related to a common resource. MERM is the cornerstone of the TaToo ontologies, meaning that all resources annotated within TaToo will contain tags according to the MERM model.

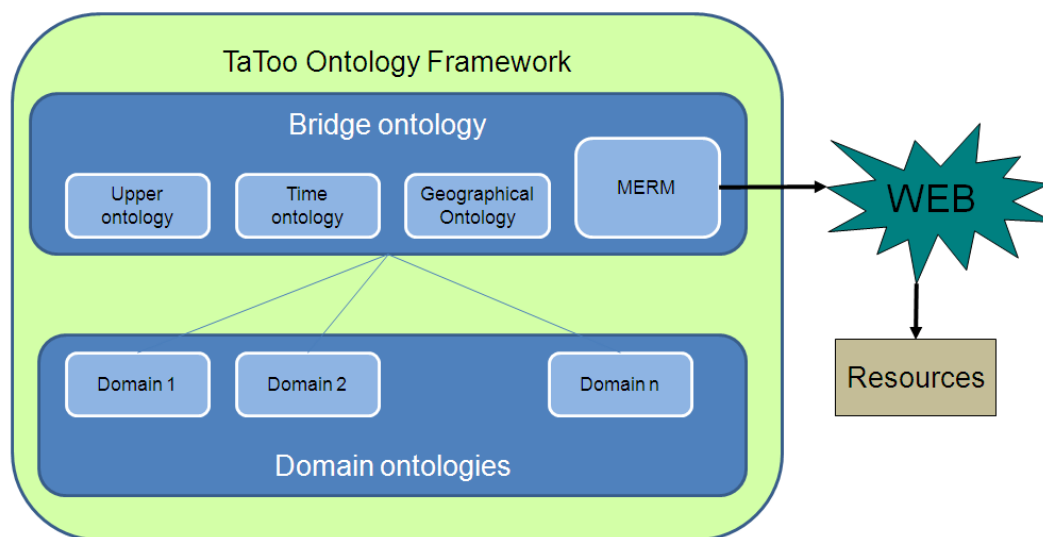


Figure 14: High-level overview of the TaToo Ontology Framework

On the other hand, MERM only impose a set of minimal structure to the TaToo annotations. However, resources must be tagged using elements from the different sub-domains within the environmental domain. The combination between formal tagging and the existence of different sub-domains implies the need of several domain ontologies. This means that the TaToo Ontology Framework must be composed of at least MERM and a set of domain ontologies. But in order to allow **cross-domain search**, the TaToo Ontology Framework should allow semantic interoperability between different

domain concepts. The approach evolved since deliverable D3.1.1, where a simple hybrid method for ontology integration was proposed, to the current approach explained in deliverable D3.1.3, where the hybrid approach is still followed in an even more lightweight fashion. This approach implies that common elements of the domain ontologies have been abstracted and mapped to a common and really simple ontology in order to allow semantic interoperability. Within TaToo this common ontology is called Bridge Ontology.

Therefore, a high level description of the TaToo Ontology Framework can be seen in Figure 14.

The ontologies defined in the Figure 14 are the following:

- **Minimal Environmental Resource Model (MERM):** MERM is the TaToo ontology that gives structure to the annotations of a given resource, hence an effort to identify a minimal model that, without limiting the expression of specific domains, acts as a reference for past and future applications in the domain. MERM specifies concepts and properties that allow the description of what is important to say about a resource common to the environmental domain. MERM is extensible in order to allow specialization on different sub-domains, if needed.
- **TaToo Bridge Ontology (BRIDGE Ontology):** The BRIDGE Ontology main objective is to map elements coming from different domains fostering easier cross-domain integration. It acts as a bridge or hub where elements of different domains are mapped together. The BRIDGE Ontology also provides a unified way to annotate entities common to most of the sub-domains of the environmental domain. In this sense, the BRIDGE Ontology contains MERM plus a set of common domain elements, such as time and geographical ontologies. Domain ontologies can map to these common elements in order to achieve cross-domain semantic interoperability.
- **Domain ontologies:** A set of ontologies from the different sub-domains that are mapped to the common elements of the BRIDGE Ontology.

In order to implement the different ontologies depicted in Figure 14 we surveyed existing work. The main conclusions for this survey are the following:

- **Existing vocabularies:** In order to foster reusability and common understanding, existing shared vocabularies have been studied to define concepts and properties of MERM and BRIDGE ontologies. As some of the most widely used vocabularies, OWL ontologies derived from Dublin Core, FOAF and SIOC are used to this purpose.
- **Upper ontologies:** The hybrid approach for semantic interoperability explained in deliverable D3.1.1 stands on the usage of a common upper ontology to map elements from different domains. Among others we surveyed the following upper ontologies as candidates to fulfil this goal:
 - SUMO: The Suggested Upper Merged Ontology. They are being used for research and applications in search, linguistics and reasoning
 - DOLCE: Descriptive Ontology for Linguistic and Cognitive Engineering
 - SWEET: The Semantic Web for Earth and Environmental Terminology (SWEET). These SWEET ontologies are perhaps the best candidates for upper ontologies in TaToo. They provide semantic descriptions of Earth system science, and are widely used in the environmental domain.

Finally, the TaToo ontology framework did not use any of this ontologies, but fragments of those (especially from SWEET) in order to create a few upper level domain concepts for the purpose of the Validation Scenarios.

- **Common shared elements on the environmental domain:** Common to the most sub-domains of the environmental domain, time and geographical locations are the most widely named elements according to the requirements and competency questions gathered from the users. There are existing ontologies and related technologies dealing with these two concepts. We have focused on three well-known ontologies: W3C Time Ontology (OWL-Time),

GeoNames Ontology (GEONAMES) and NUTS (NUTS) to provide a common way of dealing with time and location aspects in TaToo.

Figure 15 shows a detailed overview of the final TaToo ontology network:

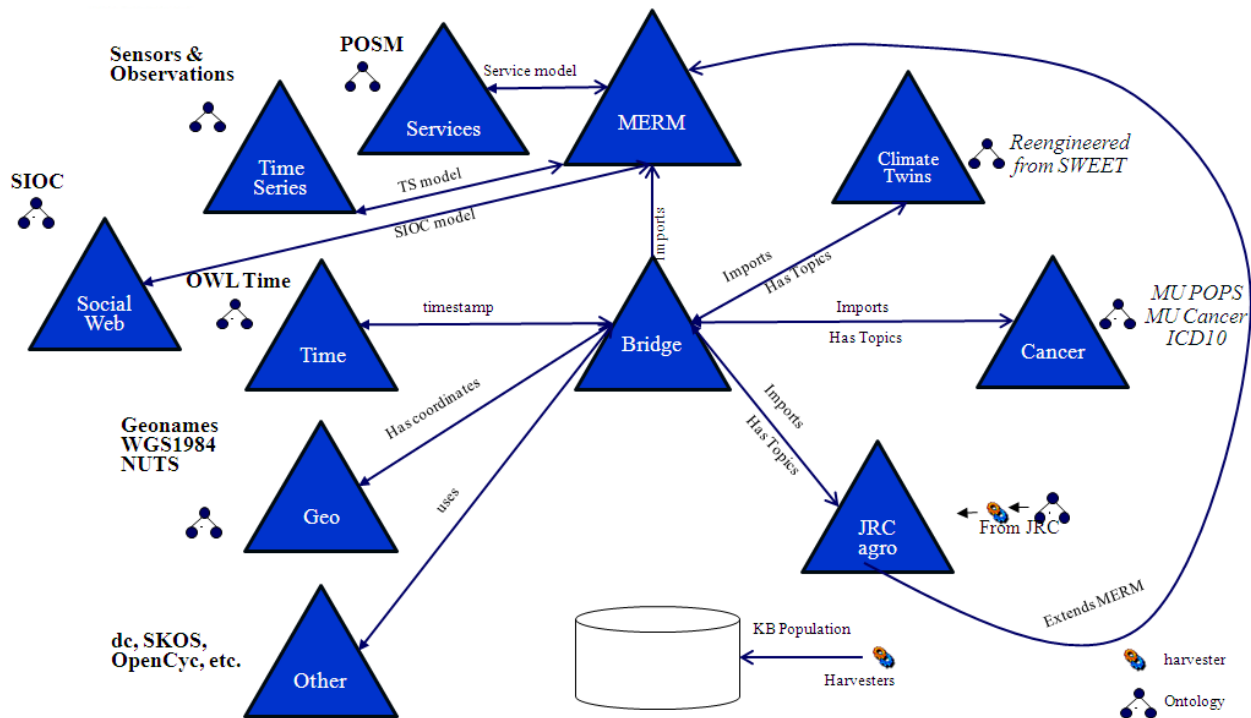


Figure 15: TaToo Ontology Network

2.3.2 Minimal Environmental Resource Model (MERM) and Bridge Ontologies

MERM has been defined above as the TaToo ontology that gives structure to the annotations of a given resource, hence an effort to identify a minimal model that, without limiting the expression of specific domains, acts as a reference for past and future applications in the domain. MERM specifies concepts and properties that allow the description of what is important to say about a resource common to the environmental domain. MERM is extensible in order to allow specialization on different sub-domains, if needed.

Therefore in the conceptualization of MERM three main classes are defined as can be seen in Figure 16. These are *Resource*, *AccessInfo*, and *Annotation*

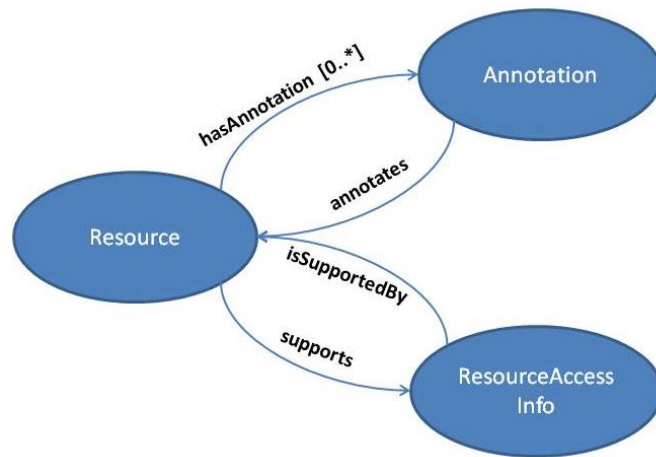


Figure 16: MERM conceptual main classes

- **Resource** is the main class for MERM. Resource should include resource management information as publisher, provider, date of creation, etc. For the first iteration, the search method implemented was based only on the *Annotation* class structure, while in v2 some resource metadata associated to the *Resource* and *ResourceAccessInfo* classes is taken also into account. The *Resource* class includes annotation management information (author, date of creation, etc.). A resource presents a set of annotations describing what the resource is about and annotations used to discover the resource.
- **ResourceAccessInfo** is about accessing the resource. We assume that a resource has a single AccessInfo. The content of the access information depends on the type of the resource. Although access information for a web document would be different from access information for a web service, in TaToo ontology we assume a single way to access to any resource; through its URI. The *ResourceAccessInfo* class has not been developed to its full potential, as it is not in the TaToo scope to actually ease the process to access to Resources.

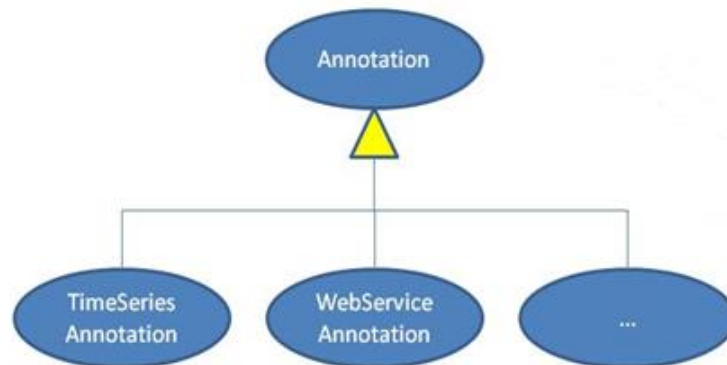


Figure 2-17: MERM Annotation conceptual model

- **Annotation** describes what a resource is about according to the people tagging the resource. It is therefore an opinion about the resource, which might or not be true. The content of an annotation depends on the type of the resource (an annotation for a web document would be different from an annotation for a web service). Saying that “some resource is related to some domain concepts” would be the most basic type of annotation in TaToo. Resources can be tagged and discovered taking annotations as basis.
- **Evaluations** are also part of the MERM model, allowing users to evaluate both resources and resource annotations according to various evaluation criteria.

- **Linking** part of the MERM model enables TaToo resources to be linked by two types of links. The first type refers to the so-called typed links defined by a selected set of well-established properties such as *rdfs:seeAlso*, *rdfs:isDefinedBy*, *dcterms:hasPart*, *dcterms:isPartOf*, etc. The second type refers to similarity links that link two similar TaToo resources. The MERM model introduces its own property (*merm:similarTo*) that describes the similarity links.

The main objective of the BRIDGE Ontology is allowing cross-domain tagging and search of resources within TaToo. In order to do that several common elements from the environmental domain have to be modelled and present in the ontology allowing mappings from the domain. Ontology links used for cross-domain include the following properties from RDFS, OWL, Dublic Core (dc), and MERM :*rdfs:seeAlso*, *rdfs:isDefinedBy* -subproperty of *rdfs:seeAlso*-, *owl:sameAs*, *owl:differentFrom*, *dc:hasPart*, *dc:isPartOf*, *dc:isReferencedBy*, *dc:isReplacedBy* and *merm:similarTo*.

For the conceptualization of the time and geographical concepts we have decided to reuse the ontologies W3C Time, GeoNames, and NUTS. Mappings from the domain ontologies to these TaToo reference ontologies are allowed to perform cross-domain search.

As an implementation convention, within TaToo all domain elements subject of annotation must be subclasses of the *bridge:Topic* class.

The BRIDGE Ontology imports MERM, plus the environmental BRIDGE elements (W3C Time, GeoNames, and NUTS ontologies).

In particular, the relation between MERM annotations and domain elements is managed using some BRIDGE Ontology elements. When tagging with domain elements, a relation between *merm:Annotation* class (or subclasses) and a domain element must be produced in order to link the resource to the domain. MERM provides a classification of types of annotations (web annotations, time series annotations, web service annotations, etc.). The BRIDGE Ontology provides a *bridge:Topic* class that allows to create the link between domain elements and the MERM annotations. In order to do that a super property (*merm:subject*) and a set of sub properties have been defined aiming at bridging MERM and domain elements as it is shown in Figure 18.

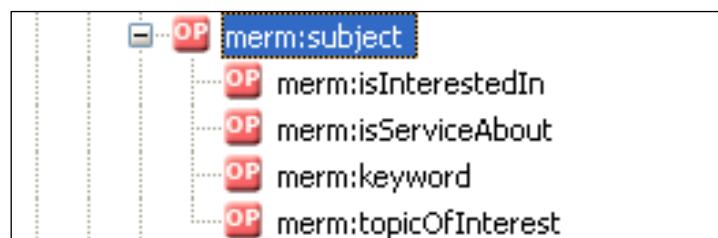


Figure 18: “Subject” Super Property

Several consequences have been obtained from those relations:

- *merm:subject* is used when we want to link some general annotation (not typifying) with *merm:Topic* class (domain elements)
- *merm:keyword*, *merm:isServiceAbout*, etc. are used to link some typified annotation with *merm:Topic*. The subPropertyOf relation in this case means that some typified annotations X also have a *merm:subject* property that holds also the relation to *bridge:Topic*.

We considered the use of *dcterm:subject* instead of creating a new super property, but it produced some problems when we imported the GeoNames Ontology and two of its properties: *GeoNames:featureClass*, and *geoname:featureCode* were also subclassified in *dcterm:subject*. Therefore “Subject” is a MERM property.

Summarizing, for each type of annotation the BRIDGE and MERM ontologies provide a specific property linking to the domain elements represented by the *bridge:Topic* class:

- As a matter of example, *merm:WebServiceAnnotation* is related to the *bridge:Topic* through the *merm:isServiceAbout* property and *merm:WebAnnotation* relates to *bridge:Topic* using the *merm:keyword* property.
- Furthermore, when the resource cannot be typified as any specific type of annotation in MERM, the annotation should be done directly using the *merm:Annotation* class. This class is linked to the *bridge:Topic* through the *merm:subject* property.

Procedure to add new ontologies to the TaToo Framework

One of the core goals of TaToo is that of allowing cross-domain discovery of resources in the environmental domain. The environmental domain is in fact composed of several sub-domains such as the ones represented in the TaToo Validation Scenarios. There is no single standard ontology or shared vocabulary that encompasses the entire environmental domain, so the TaToo ontology framework faced the challenge of being capable of dealing with annotations from different domain ontologies while enabling a certain degree of cross-domain discovery capabilities, and at the same time offering a unified, extensible and relatively simple ontology framework.

The TaToo Validation Scenarios have created several different domain ontologies to better describe their own sub-domains. The TaToo ontology framework has been designed and implemented in a way that nothing precludes in adding new domain ontologies to the network if the new ontology follows the relatively simple set of requirements and guidelines presented in this section.

To better understand the guidelines and requirements, Figure 19 shows a simplified view of the TaToo ontology network emphasizing the aspects related to the integration of domain ontologies.

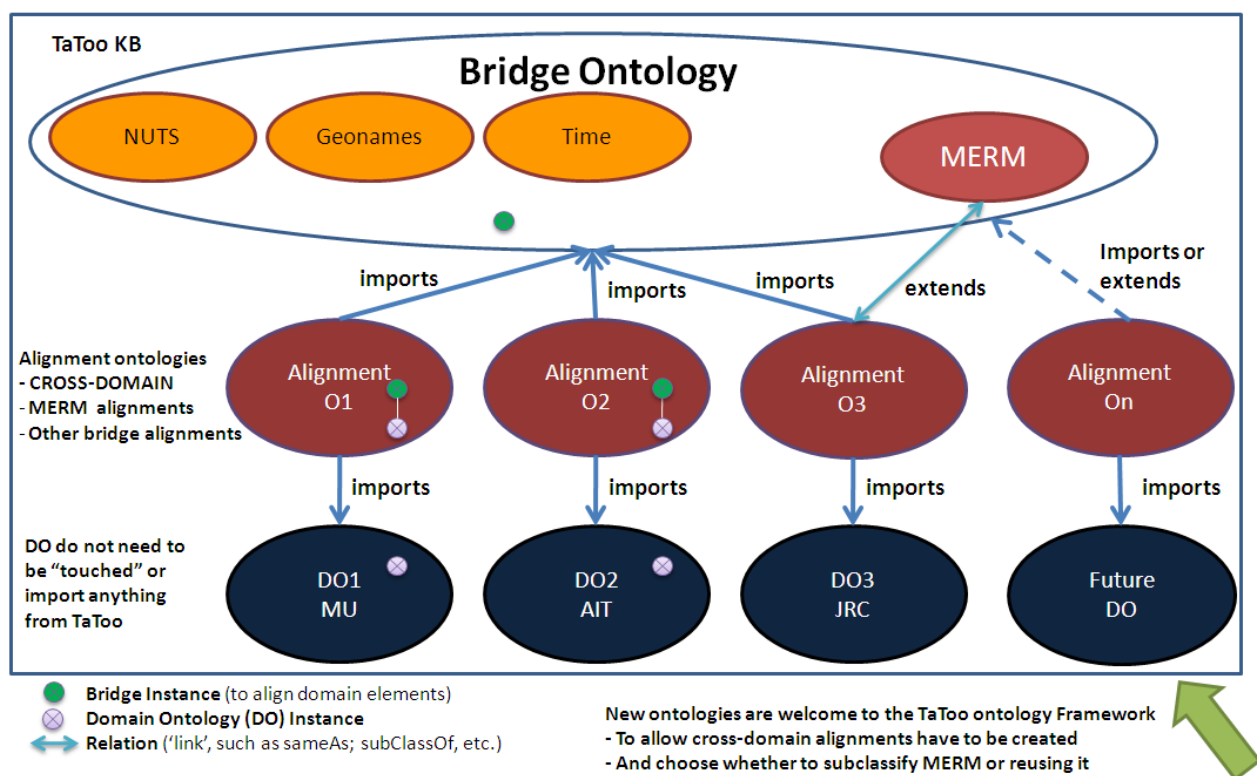


Figure 19: TaToo ontology framework to include new domains

The upper part of Figure 19 shows the main common ontologies of the TaToo Framework already described in section 0. The BRIDGE ontology is represented as an umbrella where all the rest of the supporting common ontologies (MERM, Geonames, NUTS, and W3C Time) are encompassed. On the lower part a set of domain ontologies is represented, including the ones developed for the TaToo Validation Scenarios (named as the main partner in charge of each scenario – MU, AIT, and JRC), and an extra domain ontology (named Future DO) that hypothetically would have to be integrated in

TaToo to allow annotation of resources for a new sub-domain not covered by the current ontology framework. The middle part shows a set of alignment ontologies that will contain the mappings between domain ontologies and upper level bridge elements. Note that for the implementation sake these alignment ontologies are not a necessity, as the alignments could be placed either on the domain ontology or in the bridge ontology, but they serve on the one hand to the to add new ontologies to the TaToo ontology network, and on the other hand to keep both domain ontologies and bridge ontologies untouched after the alignment process.

As explained in section 0, the main purpose of the TaToo domain ontologies is that of offering new domain elements (individuals) as topics (individuals of subclasses of *bridge:Topic*) to allow semantic annotation and search of resources following the MERM model. This implies that the focus of TaToo is not in interoperability issues, but of allowing a simple and effective way of adding new domain elements for the annotation and discovery, taking into account cross-domain considerations.

Cross-domain mappings are the most complex issue when adding a new ontology. The current TaToo ontology framework contains three main domain ontologies, but TaToo aims at a better coverage of environmental sub-domains, so new ontologies are likely to be added to the framework in the future. Therefore, allowing direct mappings between different domain ontologies could potentially end-up in a really unmanageable network of interconnected ontologies. To avoid this issue, we propose to abstract the mappings in a way that all classes or individuals subject of mapping will be available at the bridge ontology. In order to do so, the procedure to follow when a possible candidate mapping between two domains is detected is explained later among the guidelines.

The requisites for introducing a new domain ontology into TaToo are the following:

- The new domain ontology must be an OWL2 RL ontology.
- The classes of individuals that can potentially be used for annotation (topics) must be made explicitly subclasses of the *bridge:Topic* class of the bridge ontology.
- The domain ontologies must provide individuals of the topics (subclasses of *bridge:Topic*) that will be used for annotation. TaToo does not allow annotation of resources using classes in order to ease the inference process and get an optimal performance.

Therefore, the suggested guidelines for introducing a new sub-domain into TaToo can be summarized as follows:

- For the new domain ontology, create a new alignment ontology that imports both the new domain ontology and the BRIDGE ontology. As BRIDGE already contains the rest of the TaToo common ontology network, the new alignment ontology has at its disposal access to other ontologies such as MERM, GeoNames, NUTS, or W3C time.
- All changes will be produced at the alignment ontology level. Do not change anything in the Bridge Ontologies (or other imported ontologies). If changes need to be done at this level, the TaToo team must be contacted for considering the change.
- For cross-domain issues, the procedure to follow when a possible candidate mapping between two domains is detected is the following:
 - The domain ontology owner should check whether the ontology element is already available within the BRIDGE ontology. The TaToo team can be contacted for cross-checking to this extent.
 - In case the element to align is not available, a new ontology element already aligned to the other domain ontology will be produced by the TaToo team. The TaToo team will deploy the new version of the BRIDGE ontology that must be updated in all alignment ontologies that exist in the TaToo Knowledge Base (this will also be done by the TaToo team).
 - The domain ontology owner is then able to map the bridge ontology element to its own domain element using the cross-domain properties explained in section 0 within its own alignment ontology.

- The TaToo model can be also extended in case of need. MERM provides extension points especially to define new annotation types or properties not covered in the common model. JRC domain ontology is a good example of such extensions.
 - However, do not make these changes directly in MERM, but in the alignment ontology.
 - Before doing that is strongly recommended to contact with the TaToo team to ensure that the extension is really needed and sound. The TaToo team could decide to promote your extension to the general model for better cross-domain coverage, or keep it only for your domain.
- When the domain ontology is produced following these steps, the domain ontology owner will have to provide the domain ontology and the alignment ontology to the TaToo team in order to deploy both in the TaToo ontology framework. Once this is done, external users will be able to perform annotations and search based on the new topics specified.

3. Conclusions

This contribution has presented the final iteration of the implementation of the TaToo Framework.

The deliverable reported the combined work of tasks T4.1, T4.2, and T4.3 of WP4, therefore acting as glue of the overall implementation performed in WP4. This final iteration is built on top of the work performed in the previous two iterations, reported in deliverable D4.1.1 [19] and deliverable D4.1.2 [20].

This document reports on the final working version of the TaToo Framework: components and services, which can be accessed both via the TaToo Portal and the TaToo Public Services. Most of the software developed within the TaToo Framework follows an Open Source licensing schema.

The aim of the TaToo Framework is to provide an infrastructure to fill the discovery gap between environmental resources and users. The TaToo Framework enables experts as well as arbitrary users to share trusted and reliable information and to allow easy discovery and semantically enhanced tagging of existing environmental resources.

The major achievements and results of the final iteration of the TaToo Framework implementation are:

- The decision on implementation issues based on existing state-of-the-art solutions on the semantic discovery, tagging, and Linked Data fields;
- The final delivery and consolidation of the TaToo Framework Architecture implementation viewpoint;
- Consolidation and enhancement of the components and services of the TaToo Framework, including:
 - The realization of the TaToo Portal, the main front-end of the TaToo Framework, as a showcase client of the TaToo Public Services using a unified technology. The final implementation of the TaToo Portal is more user-centric and compact, taking into account the feedback from the Validation Scenarios and external users;
 - The provision of a set of public services exposing the main functionality provided by TaToo. These public services are used by the TaToo Portal components and TaToo use cases, and can potentially be used by developers of external clients to interact with the TaToo system;
 - A major work on TaToo harvesters implementing new harvester connectors and therefore gathering a substantial set of resources and annotations in the TaToo Knowledge Base for evaluation purposes.
- Major work in consolidating, improving and complete the ontologies that are at the core of the TaToo ontology framework, including cross-domain search, Linked Data extension, similarity between resources and multilingual aspects, following deliverable D3.1.3 where the basis for the TaToo ontology framework design was established;
- A detailed description of all identified components. It follows and supersedes the deliverable D3.1.3[2] where the basis for the detailed design of the TaToo Framework was established;

The appropriate licensing schema will be provided in the final exploitation deliverable, but the work done in WP4 follows an Open Source license, while frameworks used may vary their licenses if used for commercial purposes (i.e. OWLIM).

4. Acknowledgements

“The research leading to these results has received funding from the European Community's Seventh Framework Programme (FP7/2007-2013) under Grant Agreement Number 247893.”

5. Bibliography

- [1] Dihé P., Schlobinski S., TaToo Semantic Service Environment and Framework Architecture – V1, Deliverable 3.1.1 of TaToo Project, Public Document, (2010)
- [2] Dihe P., Avellino G., Petroncio L., Pariente T., Fuentes J.M., Yurtsever S., Rizzoli A., Nesic S., Božić B., Schimak G.: Semantic Service Environment and Framework Architecture V3, Deliverable 3.1.3 of TaToo Project, Public Document, (2012)
- [3] OASIS Reference Model for Service Oriented Architecture 1.0. Committee Specification 1, 2 August (2006). <http://www.oasis-open.org/committees/download.php/19679/soa-rm-cs.pdf>
- [4] Božić B., Schimak G.: Requirements Document V3, Deliverable 2.3.3 of TaToo Project, Public Document, (2011)
- [5] OASIS Reference Model for Service Oriented Architecture 1.0. Committee Specification 1, 2 August (2006). <http://www.oasis-open.org/committees/download.php/19679/soa-rm-cs.pdf>
- [6] Bloch J., Pozefsky M., JSR 10: Preferences API Specification, Java Specification Request, 09 May (2002), <http://www.jcp.org/en/jsr/detail?id=10>
- [7] Google Web Toolkit, <http://code.google.com/intl/it-IT/webtoolkit/>
- [8] Goetz B., Peierls T., Bloch J., and Bowbeer J.: Java Concurrency in Practice, Addison-Wesley Professional, (2006)
- [9] Beckett D., RDF/XML Syntax Specification (Revised). W3C Recommendation, 10 February 2004
- [10] Dan Brickley D, Guha R.V., RDF Vocabulary Description Language 1.0: RDF Schema. W3C Recommendation, 10 February (2004)
- [11] McGuinness D.L, van Harmelen F., OWL Web Ontology Language Overview. W3C Recommendation, 10 February (2004)
- [12] Dublin Core metadata initiative, <http://dublincore.org/>
- [13] The Friend of a Friend (FOAF) project, <http://www.foaf-project.org/>
- [14] The SIOC initiative (Semantically-Interlinked Online Communities), <http://www.sioc-project.org/>
- [15] Gómez-Pérez A., Motta E., Suárez-Figueroa M.C., NeOn Methodology in a Nutshell. http://www.neon-project.org/nw/NeOn_Book, (2010)
- [16] The NeOn Toolkit, <http://neon-toolkit.org/>
- [17] Protégé ontology editor, <http://protege.stanford.edu/>
- [18] Suárez-Figueroa, M. C., et al. NeOn D5.4.1. NeOn Methodology for Building Contextualized Ontology Networks. February (2008)
- [19] Fuentes J.M., Sanguino M.A., Yurtsever S., Dihe P., Nesic S., Avellino G., Petronzio L., Božić B., Schimak G.: Semantic Framework Implementation Prototype (early), Deliverable 4.1.1 of TaToo Project, Public Document, (2010)

- [20] Fuentes J.M., Sanguino M.A., Yurtsever S., Dihe P., Nesic S., Avellino G., Petronzio L., Božić B., Schimak G.: Semantic Framework Implementation Prototype (early), Deliverable 4.1.1 of TaToo Project, Public Document, (2010)

Bezpečnost sociálních sítí na Internetu

Jan Ministr

VŠB – Technická univerzita Ostrava, Ekonomická fakulta, Sokolská 33, Ostrava
jan.ministr@vsb.cz

Abstrakt

V příspěvku jsou popsány koncepty sociálních sítí v kontextu s bezpečností informací, které tyto sítě obsahují. Sociální sítě nabývají stále více na své popularitě, a tím pádem stále roste význam zabezpečení citlivých informací v nich a soukromí jejich uživatelů. Článek popisuje jednotlivá specifika bezpečnosti práce na sociálních sítích a některé zjištěné charakteristiky z hlediska bezpečnosti uživatelů v ČR.

Abstract

The paper describes the concepts of social networking in the context of information security, these networks contain. Social networks are becoming more on their popularity, and thus increasingly important security sensitive information in them, and the privacy of their users. This paper describes particularities safety on social networks and characteristics of some of the observed safety users in the Czech Republic.

Klíčová slova

Informační bezpečnost; Web 2.0; ICT; doména; profil.

Keywords

Information Safety; Web 2.0; ICT; Domain; Profile.

1. Úvod

V dnešní rozvinuté informační společnosti uživatelé Internetu na základě webových služeb Web 2.0 spoluvytvářejí a sdílejí informace novým způsobem, jehož klíčový postulát definoval (O'Reilly, 2006) jako vytváření aplikací, které se budou stále zlepšovat díky síťovému efektu s přibývajícím počtem uživatelů a přechodem od centralizovaného zpracování služeb k decentralizaci. Jednou z forem komunikace na těchto základech jsou sociální sítě na Internetu. Obecně termín sociální síť definoval sociolog J. A. Barnes již v roce 1954 jako sociální skupiny nebo komunity propojené na základě společných zájmů, které mohou být rodinné, ekonomické, politické nebo kulturní. Obecně každý jedinec je od svého narození členem určité sociální skupiny a příslušnost k určité sociální skupině jej provází celý život, což se projevuje v jeho citech, chování a myšlení.

V rámci Webu 2.0 se pod termínem sociální síť rozumí každý systém, který umožňuje vytvářet a udržovat seznam navzájem propojených kontaktů, přátel (Pavlíček, 2010). Programy, které podporují komunikaci a různé služby v oblasti sociálních sítí se nazývají socialware a umožňují na jedné straně jejich uživatelům služby jako: publikování různých informací, vkládání fotografií a alb, vytváření deníků, hledání různých známých, profesních kolegů a vytváření pracovních týmů. Na druhé straně ale největší slabinou těchto sítí je jejich bezpečnost před zneužitím informací uložených v dané sociální síti a otázka autenticity uživatelů, protože primárně uživatelé sociálních sítí tvoří hlavní obsah sociální sítě, nikoli funkcionalita sociální sítě jako taková. Toto tvrzení je podpořeno tzv. síťovým efektem platným pro média a ICT, který popsal Robert Metcalf jako užitečnost sítě, která roste se čtvercem počtu připojených uživatelů (na dané křivce ale existuje bod, od kterého pak již roste užitečnost sítě exponenciálně). Daná sociální síť je tedy pro jejího uživatele tím více užitečnější, čím více dalších uživatelů přitahuje.

2. Koncept sociální sítě jako komunity

Každá komunita ze sociálního pohledu, včetně komunity, která vznikla na sociálních sítích Internetu je založena na třech základních prvcích (Mládková, 2005), kterými jsou:

- *Doména* činnosti, nebo-li téma, která určuje:
 - Hlavní důvod pro členství a spolupráci v dané komunitě.
 - Sjednocující prvek dané komunity.
 - Způsob práce, terminologii a nástroje v závislosti na charakteru domény.
- *Mezilidské vztahy*, které jsou vytvářeny na základě:
 - Společných aktivit dané komunity a jejichž výsledkem je pocit sounáležitosti k dané komunitě, protože k práci v komunitě nelze nikoho nutit, formální uživatel bez pocitu sounáležitosti ke komunitě zůstává na jejím okraji a není pro danou komunitu přínosem.
 - Angažovanosti členů komunity, která je klíčová pro realizaci aktivit komunity, protože jinak daná komunita postrádá smysl své existence.
- *Sdílení informací spojené s tvorbou znalostí*, které je podmíněno:
 - Ochotou členů komunity vytvářet a sdílet znalosti navzájem, protože uživatel sociální sítě, který není ochoten sdílet znalosti nebo má problémy s komunikací jak technologickou, tak sociální, bývá zpravidla postupně z dané komunity vyloučen.
 - Tvorbou a následným sdílením vlastních zkušeností, nástrojů, postupů, dovedností apod.

Členové komunity na sociálních sítích mají zpravidla od dané komunity určitá očekávání ve formě přidané hodnoty a pokud se jejich očekávání nenaplní, pak z dané komunity odcházejí.

3. C4P koncept sociální sítě jako komunity

Pohled na komunitu, který je zaměřen na ICT požadavky, které musí úspěšná znalostní komunita splňovat popsal (Hoadley, 2005) pomocí modelu *C4P*, který má následující strukturu

- *Content (obsah)*, kdy by nástroje, které jsou používány danou komunitou by měly dobře podporovat jeho pohodlnou tvorbu, archivaci, vyhledávání apod.
- *Conversation (konverzace)*, kdy kvalitní nástroje pro komunikaci přímo podporují vznik kvalitního obsahu.
- *Connections (propojení)*, které je postaveno na moderních ICT.
- *Information Context (informační kontext)*
- *Purpose (účel)*, který je obdobou domény, která je popsána v předchozím pohledu.

Virtuální sociální síť na Internetu, aby byla úspěšná musí splňovat všechny předchozí body modelu *C4P*, jinak jí v současné době hrozí, že nebude potencionálními komunitami akceptována.

4. Web koncept sociální sítě jako komunity

Sociální síť je služba založená na webových technologiích, která nabízí jedincům používajícím takovou síť (Boyd, Ellison, 2007), která vlastní 3 základní možnosti, a to:

- *Vytvoření profilu uživatele* (veřejného nebo poloveřejného) v rámci této sítě.
- *Definice seznamu* dalších uživatelů v rámci této sítě, se kterými je daný jedinec *propojen*, přičemž povaha a pojmenování těchto propojení se mohou v různých sítích lišit.
- *Umožnění zobrazit a procházet seznam uživatelů* uživatelům sítě, s nimiž jsou spojeni a zároveň procházet tyto seznamy i u jiných uživatelů.

5. Specifika bezpečnosti sociálních sítí

Moderní právo je založeno na teritoriálním principu, čemuž se internet svou podstatou jako celek vymyká, protože globálnost internetu znemožňuje jeho právní specifikaci a regulaci na základě existující (teritoriální) legislativy. Změnu tohoto stavu si vynucuje vstup komerčních subjektů na internet. Právníci se snaží řešit tuto situaci v následujících oblastech:

- *Uzavírání obchodních a spotřebitelských smluv* prostřednictvím internetu zůstává otázkou dohody smluvních stran, stejně jako místní příslušnost soudů či rozhodčích orgánů

v případných sporech. Pokud není tato část právního vztahu ve smlouvě stanovena, předpokládá se využití domovských institucí zákazníka nebo spotřebitele, přičemž vymahatelnost práv vzniklých na základě domácího právního systému vůči zahraničním subjektům je ovšem v některých případech velmi malá. S touto oblastí těsně souvisí jednoznačné určení okamžiku jednání v čase, protože veškeré dokumenty, které nejsou označeny časovým razítkem, jsou snadno manipulovatelné.

- *Trestněprávní oblast* není dořešena. Většina států vlastní právní systém, ve kterém je jasně popsáno co je považováno za trestné, procesní stránka věci je mnohdy zcela odlišná. V praxi v rámci evropského právního prostoru je sice možné vydat občana do jiného státu, pokud existuje podezření, že tam spáchal čin, který je v daném státě trestný, ale zatím se tato možnost příliš často nevyužívá. V případě střetu s jinou kulturou je pravděpodobnost úspěšného řešení sporu právní cestou ještě podstatně nižší.
- *Autorskoprávní vztahy* se snaží producenti namísto nového obchodního modelu marně chránit dodatečným vkládáním bariér jako:
 - Šifrování obsahu zabezpečující nelegální kopírování.
 - Znalost unikátních hesel zabezpečující SW aplikace a HW omezení v přenosovém řetězci.
 - Dobrovolná identifikace koncového zákazníka až do podoby digitálního podpisu a certifikace.
 - Propojení fyzické osoby s daným HW zařízením, u kterého končí jednoznačně identifikovatelná a zaznamenaná stopa činnosti (IP adresa, ale jednoznačně neurčuje konkrétního uživatele, který stiskl dané tlačítko). V některých zemích ale lze vysledovat legislativní tendence k přesunutí určité právní odpovědnosti i na majitele počítače s danou IP adresou v daném časovém období, stejně jako poskytovatele připojení - a to jak klasické ISP, tak provozovatele domácích wi-fi sítí.
- *Licenční podmínky* se provozovatelé aplikací a služeb Webu 2.0 řeší pomocí a priori stanovených „podmínek užití“ dané služby, které musí každý uživatel odsouhlasit během registrace, nebo s nimi obvykle vyjádří souhlas tacitně - pouhým užitím dané služby (pokud není registrace vyžadována). Tyto podmínky užití většinou obsahují licenci uživateli k použití dané aplikace, licenci provozovateli k nakládání s daty uživatele, práva a povinnosti jednotlivých smluvních stran, výhrady zodpovědnosti a další „informativní“ položky. V praxi je však právní gramotnost uživatelů internetu poměrně nízká a lze předpokládat, že většina uživatelů podmínky použití při registraci nestuduje nijak podrobně (nebo je nečte vůbec, i když většina provozovatelů si někdy vymínuje možnost změny podmínek užití s okamžitou platností a bez nutnosti opětovného schválení uživatelem - či upozornění).
- *Ochrana osobních údajů* je základním nástrojem ochrany soukromí osob, kdy je vyšší váha vždy přisuzována zachování soukromí (směrnice č. 95/46/EC – EU Privacy Directive). Porušení tohoto principu je nezvratné (jednou zveřejněný údaj je velmi obtížné znovu utajit). Bohužel v současné době je možné pozorovat celkový odklon od striktní ochrany soukromí ve jménu tzv. všeobecně prospěšných aktivit typu „boje proti terorismu“ či politických zdůvodnění proklamujících, že „slušný člověk nemá co skrývat“. V ČR je ochrana osobních údajů implementována v podobě zákona o ochraně osobních údajů (zák. č. 101/2000 Sb.), která zavádí povinnou registraci správců osobních údajů stejně jako povinné náležitosti takového zpracování (zahrnující ukládání a přenos) týkající se bezpečnosti, protokolování apod. Provozovatel komunitního portálu musí zjistit, zda není povinným subjektem dle zákona o ochraně osobních údajů - správcem či zpracovatelem osobních údajů a registrovat se, pokud pro něj neplatí výjimka z této oznamovací povinnosti dle § 18 zákona o ochraně osobních údajů.
- *Bezpečnost informací (uložení informací a jejich toky)* je považována za bezpečnou, pokud je adekvátně zajištěna jejich:
 - důvěrnost,
 - integrita,
 - dostupnost,
 - nepopiratelnost akcí,

- *autenticita (ověření zdroje původu),*
- *zajištění soukromí (případně anonymity).*

U komunitního portálu, kde nedochází k uzavírání právních vztahů, je otázka informační bezpečnosti realizována v zásadě pouze na oblast zajištění důvěrnosti citlivých údajů a oblast zajištění soukromí. Pokud jsou některá data považována za citlivá (jako např. osobní údaje), je úkolem provozovatele komunitního portálu zabezpečit uložení a přenosy těchto dat neveřejným způsobem. V případě, že se jedná o informace spojené s konkrétní osobou, lze je z pozice provozovatele sdělovat pouze tak, aby konkrétní osoba nemohla být identifikována (Pavlíček, 2010). Při zajištění bezpečnosti je nutné vzít v úvahu lidský rozměr problému, kdy proces zpracování informací závislý i na chování lidí. Proto by měl v dané komunitě existovat systém pravidel, jak zacházet se sdělovanými informacemi (uživatelská jména, přístupová hesla apod.) a do jaké míry důvěřovat ostatním uživatelům. Zjednodušování těchto vztahů, které je implementováno do sociálních sítí představuje významnou bezpečnostní hrozbu, protože jednotliví uživatelé vnímají důvěrnost této vazby různě - někteří se v rámci sociální sítě přátelí s kýmkoli, někteří se známými, jiní pouze se skutečnými přáteli. Takto osobně vnímaná důvěrnost vazeb deformuje vztahy v sociálních sítích z hlediska bezpečnosti a ochrany soukromí. Jednotliví uživatelé přiřazují získaným informacím právě takovou důvěrnost, jak celkově vnímají danou důvěrnost daného vztahu (vazby). V komunikaci many-to-many na sociálních sítích jsou informace automaticky distribuovány všem. Změna míry důvěrnosti informace mezi jejím poskytovatelem a příjemcem může mít za následek úniky citlivých informací či dezinterpretaci sdělení.

- *Zajištění soukromí*, které je zaměřeno na ochranu proti nakládání s identitou uživatele (včetně jeho dat) bez výslovného souhlasu daného uživatele. Základním nástrojem zajištění soukromí je vyloučení identifikace subjektu (tj. anonymita, pseudonymita, ochrana osobních údajů). Každý uživatel by si měl být vědom dosahu a dopadu zveřejnění údajů o své osobě. Stejně tak by si měli uživatelé uvědomovat, že svou publikační činností více či méně ohrožují soukromí nejen své, ale i osob, o kterých se zmiňují. Je zcela na rozhodnutí provozovatele sociální sítě, zda zavede pravidla regulující publikování citlivých údajů (vč. specifikace takovýchto údajů) a do jaké míry bude regulaci provádět pomocí:
 - *Zavedení sankcí* při porušení zavedených pravidel včetně zákazu dalších publikačních aktivit
 - *Moderování příspěvků* před jejich zveřejněním, kdy jsou příspěvky uživatelů uloženy v systému provozovatele sociální sítě a je s nimi seznámen moderátor. Moderování příspěvků slouží jako nástroj cenzury příspěvků obsahujících vulgární, rasistický či obdobným způsobem nepřijatelný obsah
 - *Použití automatických systémů* založených na filtrování příspěvků s použitím seznamů zakázaných slov.
 - *Otevřenost systémů provozujícího sociální síť* vychází z podstaty principu práce komunitního serveru, který je založen na službách Web 2.0.

Otevřenost sociálních sítí můžeme hodnotit z pohledu:

- *Uživatele k sociální síti*, který lze definovat jako míru bariér pro užití dané služby (přístupová práva k různým částem systému, nucené registrace, vyžadování osobních údajů výměnou za užití, nepřehledné uživatelské rozhraní).
- *Sociální síť k uživateli* je důležitá otevřenost ve smyslu poskytování kontroly nad daty publikovanými uživatelem i samotnou funkčností služby (uživatelskou konfigurací, rozšiřováním prostřednictvím komponent dalších poskytovatelů atp.).
- *Účel sociální sítě*, který závisí na tom, zda síť byla vytvořena jako komerční (zisková) nebo nezisková. Sociální síť zaměřená na generování ekonomického zisku prostřednictvím reklamy či zpoplatnění pouze části služeb (tzv. prémiového obsahu) jsou zpravidla uživatelům zcela otevřené. Platební modely v takových službách se různí od využití paušálních poplatků, přes odstupňované „verze“ přístupu, po platby za jednotlivé akce v systému. V neziskovém sektoru sociální síť vytváří poměrně malá skupina osob a v případě úspěchu je sociální síť zpravidla zpřístupněna

veřejnosti. Neziskové komunity vzniklé okolo ústředního tématu jsou často poměrně uzavřené a noví členové si musí členství nějakým způsobem zasloužit či vybojovat. To je hlavním důvodem delší existence nekomerční sociální sítě, oproti komerční sociální síti, ve které je idea (účel) uměle konstruována a uživatelům vnucena zvnějšku.

- *Překrývání sociálních sítí* v rámci jedné virtuální sítě (např. profesionální a soukromá síť kontaktů). Určení uživatelé mohou pak jednu sociální síť využívat pro různé účely, jeden k vytváření kontaktů na své osobní přátele ve svém profilu a jiný k navazování profesních kontaktů (v postatě jde pak o dvě sociální sítě). Vzájemné propojení takových sítí vytváří napětí, které může vést k nechtěnému rozpadu některých vazeb.

6. Bezpečnost na sociálních sítích v ČR

Na základě hrubých odhadů InternetWorldStars z celkového počtu Čechů (10 190 213) činí počet Čechů připojených k internetu (6 680 800). V tabulce 1 jsou uvedeny přibližné počty Čechů připojených k nejrozšířenějším sociálním sítím v Evropě. Češi jako takoví patří mezi nejnáruživější uživatele sociálních sítí v Evropě, podle informací serveru cnews.cz „Čtyři z pěti dotazovaných uživatelů sociálních sítí kontrolují svůj status minimálně jednou denně, z toho 5 % častěji než jednou za hodinu, 35 % třikrát až čtyřikrát denně a 41 % jednou denně. V intenzitě přístupu na sociální sítě jsme v Evropě šampioni - říká český PR manažer Intelu, Pavel Svoboda.“ Vzhledem k této situaci je mimořádně důležité, aby si uživatelé sociálních sítí uvědomovali rizika spojená s tímto druhem komunikace. Bezpečnost práce na sociálních sítích mapovalo několik bakalářských a diplomových prací, které autor vedl.

Tabulka 1: Odhadované počty Čechů na sociálních sítích v roce 2011

[Zdroj: <http://www.cnews.cz/cesko-pry-propadlo-socialnim-sitim-kde-zije-nejvice>]

	Počet uživatelů	Penetrace v Česku	Penetrace mezi připojenými Čechy
Facebook	3 333 020	32,71 %	49,89 %
Google+	43 334	0,43 %	0,65 %
LinkedIn	181 703	1,78 %	2,72 %
Twitter	40 196	0,39 %	0,60 %

Poznatky získané na základě dotazníkových šetření lze z těchto prací lze shrnout do následujících klíčových tvrzení, že zhruba:

- 60% uživatelů sociálních sítí si nikdy nezměnilo své přihlašovací heslo.
- 65% uživatelů sociálních sítí nebylo obeznámeno s problematikou bezpečnosti práce na Internetu na školách.
- 20% uživatelů sociálních sítí sdílí své osobní údaje, čímž tvoří potenciál pro krádež identity.
- 30% uživatelů sociálních sítí si myslí, že nepoužívají sociální sítě bezpečně.
- 40% uživatelů sociálních sítí tvoří věková skupina 18-26 let.
- 20% uživatelů sociálních sítí tvoří věková skupina mladší 18 let a tato skupina o sobě poskytuje nejvíce citlivá data.

Vzhledem ke skutečnosti, že popularita sociálních sítí má rostoucí trend, lze konstatovat, že využívání sociálních sítí není nebezpečné při dodržování základních pravidel, kterými jsou:

- Při registraci do sociální sítě věnovat pozornost uvedeným podmínkám použití a zásadám ochrany osobních údajů.
- Používat pseudonym místo skutečného jména, aby jenom skuteční přátelé věděli, kdo tuto přezdívku používá.
- Svůj profil zpřístupnit pouze tomu, kdo je skutečným (označeným) přítelem, abychom ochránili své soukromí.
- Mezi kontakty přidávat pouze osoby, které známe osobně, chráníme tak i své skutečné přátele.

Při nedodržování výše uvedených zásad se mohou sociální sítě stát skutečně nebezpečným nástrojem zneužití v nich uchovávaných informací.

7. Závěr

Výše popsany nástroj Business Intelligence hodnotí sentiment na základě statistických funkcí, přičemž výsledek silně závisí na skutečnosti, jak kvalitně je systém naučen, tj. jakým způsobem byla vytvořena znalostní báze. Skutečnost, že je znalostní báze vytvářena vždy znovu samostatně pro každou doménovou oblast (zpravidla téma diskuse) i pro typ hledaného sentimentu sice zvyšuje nároky na tvorbu systému, ale zvyšuje úspěšnost hodnocení sentimentu zpracovaných nestrukturovaných textů.

8. Poděkování

Príspevok vznikl s podporou projektu CZ.1.07/2.4.00/17.0004 - Informace a konkurenceschopnost - INFOKON.

9. Literatura

- [1] Boyd, D. M., Ellison, N. B. (2007) Social network sites: Definition, history, and sholarship. *Journal of Computer-Mediated Communication*. 2007. vol. 13, no. 1, pp. 210-230. ISSN: 1083-6101.
- [2] Hoadley, C. M., Kilner P. G. (2005) Using technology to transform communities of practice into knowledge building communities. *ACM SIGGROUP Bulletin - Special issue on online learning communities*, 2005, vol. 25, no. 1, pp. 31-45.
- [3] Ministr, J., Ráček J., Toth, D. (2012) Visualization of the Discussion Content from the Internet. In: *IDIMT-2012, ICT for Complex Systems - 20th Interdisciplinary Information Management Talks*. Linz: Trauner, 2012, pp 297-304. ISBN: 978-3-99033-022-7.
- [4] Mládková, L. (2005) *Management znalostí*. Praha: Oeconomica, 2005. 191 s. ISBN 80-245-0878-8
- [5] O'Reilly, T. (2006) *Web 2.0 Compact Definition: Trying Again*. [on-line] <http://radar.oreilly.com-/archives/2006/12/web_20_compact.html>
- [6] Pavlíček A. (2010) *Nová média a sociální sítě*. Praha: Oeconomica, 2010. 181 s. ISBN 978-80-245-1742-1.
- [7] Škrabálek, J., Kunc, P., Pitner, T. (2012) Inner Architecture of a Social Networking System. In: *Proceedings SOFSEM'12: Trends in Theory and Practice of Computer Science*. Berlin: Springer-Verlag, 2012, pp. 530-541. ISBN: 978-3-642-27659-0.

Scaling CEP to Infinity

Filip Nguyen, Tomáš Pitner

Masaryk University, Faculty of Informatics, Lab Software Architectures and IS
Botanická 68a, 602 00 Brno, Czech Republic
{fnguyen,tomp}@fi.muni.cz
<http://lasaris.fi.muni.cz>

Abstract. Scaling CEP applications is inherently problematic. In this paper we introduce solution for scaling CEP applications that is fully distributed and aspires to scale CEP to the limits of current hardware. Our solution simplifies existent Event Processing Network abstraction and adds features on the level of CEP that change direction of its usage.

Abstrakt. Změna měřítka CEP aplikací je neodmyslitelně problematická. V tomto článku představíme řešení pro škálování CEP aplikací, které je plně distribuované a usiluje o měřítko CEP na hranici současného hardwaru. Naše řešení zjednodušuje existující abstrakci Event Processing Network a přidává další funkce na úrovni CEP, které mění směr jeho použití.

Keywords: CEP, scalability

Keywords: CEP, škálovatelnost

1. Introduction

Complex Event Processing (CEP) brings real-time processing of massive amounts of events with aggregation capabilities. This aggregation enables to correlate between events in real-time over large sliding windows. This is crucial capability of today high-end CEP technology and theory. Idea of correlating large amount of events in single window raises interesting questions about scalability of any implementation of such idea. Is it possible to scale such processing? Do it in distributed fashion? These questions are hard to answer. Our results show that move from centralized CEP to distributed model requires changes not only in technology. The move would require to evaluate basic CEP assumptions and workflow of the developer who leverages CEP.

Fig. 1 illustrates simple CEP application. One might argue that CEP application uses so called Event Processing Network (EPN) thus is distributed in some sense. This is not true, because EPN behaves towards producers as monolithic engine and its behavior is static.

In this paper we revisit the problem of scaling context aware CEP. Our approach is based on creating peer-to-peer model of processing agents without centralized coordination.

2. Complex Event Processing

Complex Event Processing (CEP) is computing abstraction for working with events and streams of events. There is a big body of literature about the topic [13][7]. The research areas in field of CEP is very wide and encompass theoretical, business oriented, technology oriented topics. In our research we are concerned with scalable pattern matching over event streams, which is a specific area of CEP. We believe this area is also the most important in whole CEP. Simpler scenarios that are concerned with filtering of stateless channels with single events can be relatively easy handled in proprietary fashion. In these simple cases CEP helps mostly by giving common terminology.

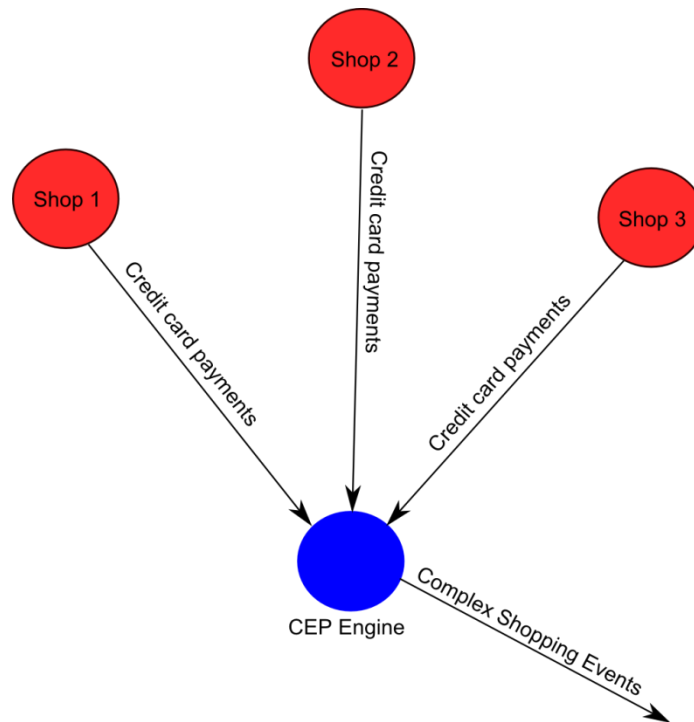


Fig. 20. Simple CEP application

It is appropriate to begin with simple example, using terminology that is mostly guided by intuition. We want to give basic understanding of CEP concepts. Simple example of pattern matching over event streams is depicted in Fig 1. Red circles represent the *event producers* and blue circle represents *CEP engine*. The event producers in our example are retail business shops that produce events representing *payment by credit card*. We view any such possible event e as variable of the first order logic. Accompanying predicate symbol $twindow(x, y)$ is designating two events x, y that happen in same time window (time window size etc. is part of the query). Function $cardnum(e)$ is returning card number and $shop(e)$ has information about the shop in which the purchase was done.

The events are entering CEP engine. The engine is able to detect specific relation between events. Intuitively but more precisely: it can select subset of events such that for the set, following logic sentence is true $\forall x, y (twindow(x, y) \wedge cardnum(x) = cardnum(y) \wedge \neg (shop(x) = shop(y)))$. In example on Fig. 1 such pattern will be selecting payments x, y which are done by 1 credit card with card number $cardnum(x)$ in different stores. Whenever such subset is identified on the stream of data, new events are generated (*Complex Shopping Events*). This concrete pattern matching have clear business oriented motivations:

- stores for A and B might be in geographical proximity
- stores for A and B might have common products for specific type of customer

To extract this kind of information is relatively easy even with existing tools (relational databases, data mining methods). So why CEP? What CEP brings is stream processing. The data are never stored, they are processed "on the go". IN terms of first-order logic pattern matching model we say that the latest subsets matched by pattern is considered in current moment. Lateness of the set is defined by the timestamp of the latest event in that set. Expressing patterns in first order logic formula frees ourselves in expression power, but for CEP there are (similarly to SQL) vast amount of language dialects that declaratively define the event patterns. All CEP dialects have one thing in common. They incorporate time of the event occurrence as first class citizen, together timewindow abstraction and with many operators that help understanding the time correlation between events, e.g.:

- event A happened before event B
- event A happened while B

Some CEP dialects even include spatial pattern recognition.

Real time access to processed data is one of the main advantages for Complex Event Processing [10]. This advantage is even strengthened if the data are mostly relevant in current moment.

2.1 Technological Tools

CEP is enabled by many software tools today. The main Open Source tool is Esper for Java and .NET. Esper is centralized CEP engine. Main advantage of it is object oriented nature and embedability into Java process. The patterns in Esper takes form of SQL-like declarative rules that are given to the engine in the form of uncompiled String, e.g.:

- `String epl = "select avg(price) from OrderEvent.win:time(30 sec)";`
- `EPStatement statement = epService.getEPAdministrator().createEPL(epl);`

IBM Active Middleware (AMiT) is another successful software suite used for CEP [21]. AMiT architecture is closer to theoretical model in current CEP literature. It eclipse based IDE with possibility to construct Event Processing Network. Event patterns in AMiT takes form of logical formulae similar to this:

- `If transaction.type="cash_check" and`
- `transaction.amount>=transaction.parameter_check_threshold`

2.2 Scaling Complex Event Processing

Scaling of CEP is regarded as nonfunctional property of CEP. We are concerned with horizontal scalability in the following categories:

- Volume of processed events
- Quantity of agents
- Quantity of producers
- Quantity of context partitions
- Availability

The volume of processed events is most important criterion for sliding window correlation. With any complex pattern, it is immediately clear that CEP engine has to correlate between all the events in the time window thus have them in the memory.

Scaling in quantity of agents is naturally given by our method because our assumption is that event producers ultimately become CEP engines.

Quantity of producers is also scaled for free by our method, because we leverage peer-to-peer model where each new producer is regarded as yet another building block of CEP application.

Quantity of context partitions is tricky scaling criterion. However, after careful look, it lies in the heart of our method. Partitioning contexts of events is the key to distribute load to infinity.

Current approaches to scalability can be categorized as follows:

1. Optimizations related to EPA assignment
2. Optimizations related to coding of specific EPA
3. Optimizations related to execution process

We believe the (1) is most important and promising area of getting true horizontal scalability for CEP applications. In this area several approaches already took place:

- Stratification
- Peer-to-peer scaling
- Vertical scaling ([2])

Current approaches to scale Complex Event Processing using stratification were studied in [6]. By stratifying the CEP architecture, it is possible to distribute load of events among more than one engine. This stratification is static way of scaling CEP. The input Event Processing Network (EPN) is stratified by algorithm into so called strata. This method mainly benefits from the fact that some event

processing agents work as filters thus are independent of the context and other agents (which are put into different strata). High throughput is achieved using this method, but it is limited by static nature.

Another way to distribute CEP is to push event processing to producers which was already studied in [9], [16]. We believe this to be very promising direction. In [4] distributed query engine is motivated by peer-to-peer file sharing. Authors identify peer-to-peer approach to be very well suited for monitoring scenarios and intrusion detection. Relaxation of design principles is identified as main way to bring good distribution of query engine.

In [5] authors identify self-organization capabilities as central requirement for large-scale distributed applications. However, CEP theory and even the implementation doesn't encompass this ideas.

3. Research Problem

CEP engines help to handle massive amounts of events. However, the inherent limitation of the idea behind pattern matching are that correlation between events must happen in centralized fashion. This paradox leads to discussion about how to scale CEP. In our example, as number of stores connected to our engine grows, one might deploy more intelligent approaches to handle number of events. In this paper we are concerned how to scale CEP to the limit of existing hardware.

In previous section we have identified current approaches to scaling. The problem with current approaches is that they still understand centralization of CEP processing and use static approaches to grow query engines. Some work has been done in peer-to-peer distribution and theories resemble existing NoSQL databases with limited context pattern matching and stream processing capabilities.

We focus on creating system that scales dynamically with the quantity of producers, context partitions, while giving guarantee of availability. Our research neglects any optimizations in vertical fashion (stratification, anything related to current communication stack or process execution optimizations).

In next section we introduce Collaborative Fully Distributed Complex Event Processing (CFD-CEP) which is our solution to this problem

4. Collaborative Fully Distributed Complex Event Processing

In this section we introduce our model for CEP. Our method is fully distributed. That means each node has exactly same semantics and is no different from any other node. The Fig 2. shows difference between typical CEP engine and our method. On the left, typical engine correlates between events from all the producers. The engine can be part of EPN but that partitions contexts (we revisit this limitation in next section). Our solution on the right is free of any EPN abstraction. Each producer is simultaneously a CEP engine.

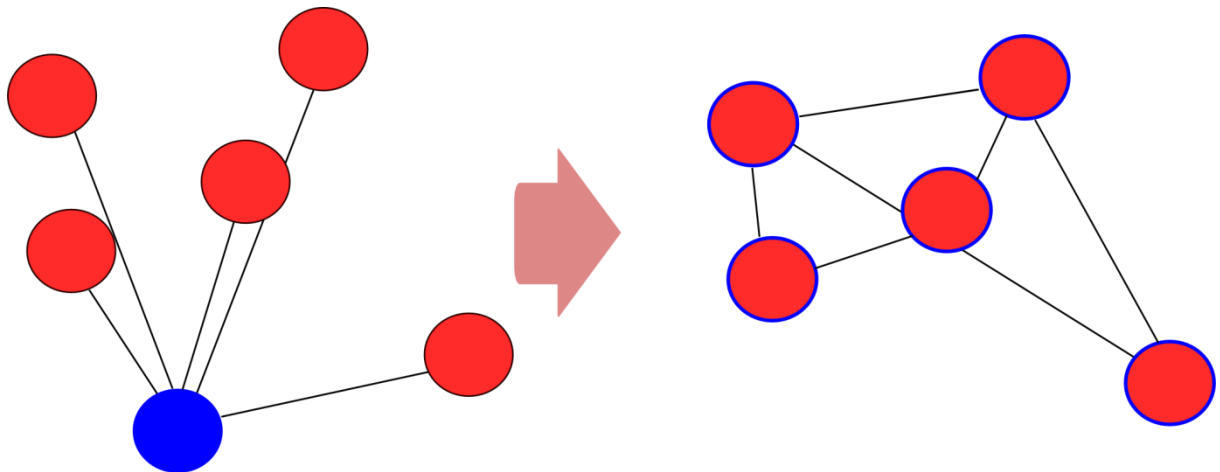


Fig. 21. The difference between typical CEP engine and our method

Our CFD-CEP model becomes distributed application in classical sense with dynamic properties and asynchronous communication channels. To achieve centralized cooperation for specific event processing tasks we use *leader election* distributed algorithms that dynamically select coordinator of specific task. Each instance of CEP engine has two main differences to typical CEP application.

1. Implementation
2. Pattern reaction syntax

Implementation has to support creation of new CEP nodes directly on producers. This alone is non-trivial requirement and by far most important problem to tackle.

Another problem of the implementation, application of distributed algorithms, is not that complex in the light of the fact that communication links are established between the nodes.

Pattern reaction syntax in current systems can generate higher level event. We introduce new events that will be used to alter structure of the graph that makes up the distributed CEP engine:

- *Node event creates/deletes processing node* (engine) or adds new event pattern to specific engine
- *Link event adds/deletes link* between two nodes

These two events grant possibility to manipulate the distributed CEP graph. Beside these, we allow to insert a new producer/engine into a graph manually.

4.1. Collaborative Subset Identification

This section describes very important problem that arises with sliding windows. Suppose we have to match following, very simple, pattern $\forall x, y (x = y)$ in environment that is depicted on Fig 3. Producers 1-5 are generating letters from alphabet. To do such pattern matching, it is necessary that CEP engine holds all the events that are in specific time window in memory and does matching among them. As number of producers grow, more computation power is needed to enable the matching.

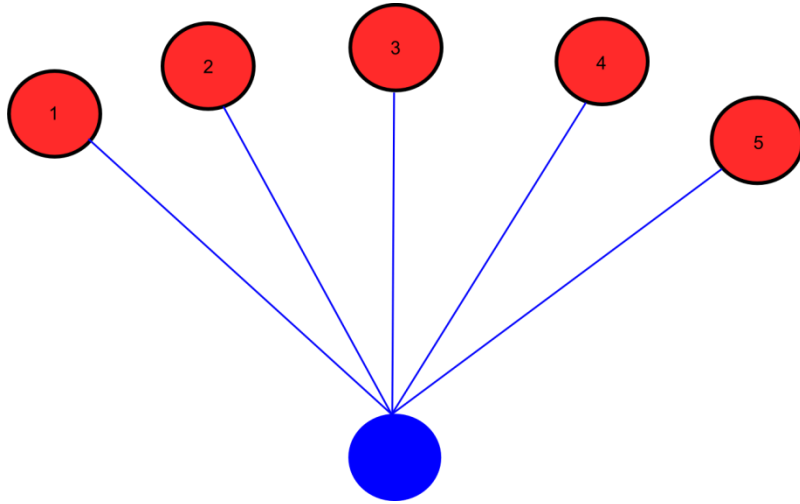


Fig. 22. Producers 1-5 are generating letters from alphabet

We argue, that distribution of letters from different sources is different. Suppose that producers 1 and 3 have higher probability of generating letter *A* than any other producers. Knowing that statically, one might add yet another engine to the application just for this simple case as depicted on Fig 4. New engine *AE* has very simple semantics. It only detects *A* events over producers 1, 3 and correlates between them. Because we knew that 1,3 have higher probability of generating such events it is also more probable that the *AE* is useful.

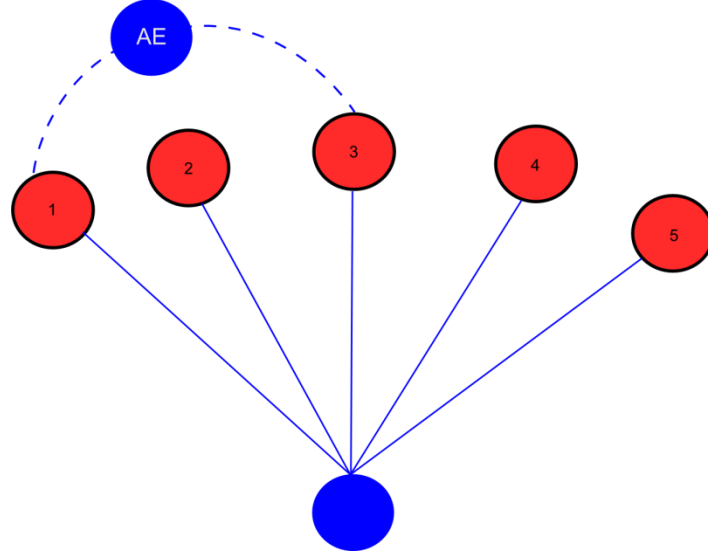


Fig. 23. Adding AE we have created another context for pattern matching

By adding *AE* we have created another context for pattern matching. It is significantly smaller than the context of the main CEP engine in the example. We could now delete links between 1, 3 and main CEP engine. That would unload event burden on main engine, but simultaneously we would decrease probability that some events that were sent by 1 and 3 would be correlated with other engines. Nevertheless, we think that this tradeoff is necessary to scale such applications.

Interesting fact to note about matching $\forall x, y (x = y)$ is, that theoretically it would be possible to match such rule, without loss of any information by putting engines between every binary combination of producers as depicted on Fig 5. This is obviously another extreme as opposed to previous solution. Number of such *AE* engines grow exponentially - $\binom{n}{2}$, where n is a number of producers. Clearly, the tradeoff between number of *AE* engines and information extraction capabilities of whole CFD-CEP engine has to be made. To make this tradeoff it is necessary to identify subsets of producers that are somehow related and create context above each of such subset for pattern matching. We call this problem *Collaborative Subset Identification Problem* and our solution is CFD-CEP. Identification of specific subsets is made by CEP patterns and context creation is done by generating *Node event* or *Link event*.

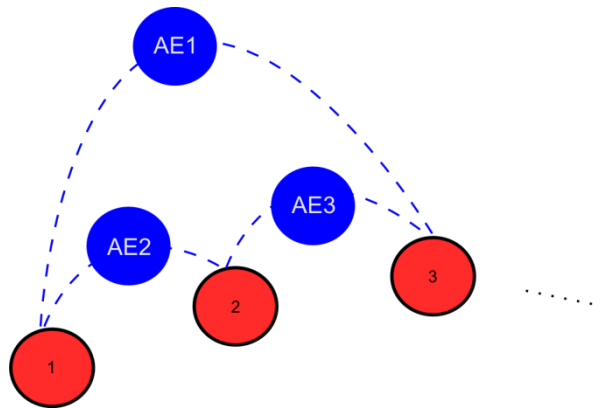


Fig. 24. Identification of specific subsets is made by CEP patterns

4.2. Evaluation Metrics

This subsection discusses possible metrics that will be part of experiment to measure performance of CFD-CEP. In distributed computing *message complexity* is known metric. This will, for sure, be used to complexity of context creation over identified subsets.

Most important measurement will be concerned with information loss relatively to performance gain of the whole system. This metric will therefore measure tradeoff situation and will prove that tradeoff has been made on reasonable basis.

4.3. Challenges

There are few, hard to solve, challenges that we face with CFD-CEP:

1. How to measure uncertainty in CFD-CEP
2. How to map statistical methods to CEP patterns to solve *Collaborative Subset Identification*

Uncertainty in CFD-CEP is given by the fact that information is lost. Measuring such loss in controlled experiment will be very demanding on high quality data set.

To solve *Collaborative Subset Identification* it is necessary to create clear guidelines how to find out that some producers are related to each other. Initial idea is to correlated coarse grained events on central CEP engine (selected by *leader election* distributed algorithm) and afterwards deploy new contexts between identified engines.

4.4. CFD-CEP Implementation

For experimental evaluation it is necessary to create prototype implementation. Implementation consists of daemon written in Java programming language. The daemon is runnable on any device with Java Virtual Machine. It is a node in the distributed application graph with knowledge only about its neighbors.

To monitor and run experiments we develop web portal for running historical data that simulate event streams. We send them into the distributed network from the web application. This web application is absolutely independent of the idea itself and acts as monitoring tool for the experiments.

We do not plan to write the daemon from scratch. After evaluation, we have selected two existing software packages/projects to base our daemon upon:

- Smartfrog
- Esper

Smartfrog is Java project developed at HP labs. It is devoted to deploy distributed applications. We leverage some of its abilities to deploy distributed system and ensure communication links between nodes. Smartfrog was intended to mostly static deployments. Our use case is highly dynamic. That might be possible obstacle, but benefits from object oriented nature of internal Smartfrog API is convincing fact to use it.

Esper is Java based CEP engine that features SQL like syntax for patterns. Together with Smartfrog we will integrate it into the daemon and distribute such package to every node in the distributed system. Thanks to lightweight implementation of Esper and possibility to embed Esper into JVM it is perfect candidate for such implementation.

We believe that our approach can be further extended to fully functional software suite, because foundations of our solution will be based on industry-proven software packages.

5. Conclusion

In this paper we have presented novel approach to CEP scaling. Collaborative Fully Distributed Complex Event Processing (CFD-CEP) allows to scale CEP applications on peer-to-peer basis. This way, CFD-CEP solution is seen as a network of identical EPAs. We further simplify notation to describe EP networks known from the literature, because we believe pattern matching over sliding windows to be the most important area of CEP research. On the other hand we add more focus on dynamic distributed aspects of EP networks and we enrich the CEP description with rigorous descriptions of CEP pattern matching rules. Rules are given in formal way, but not only that, they also feature new construct to enable distributed processing to be captured.

Because EPA is close to the producer we also achieve effective vertical scaling regardless of computing architecture being used. While experimental evaluation is still pending, we have presented design of target system and identified possible implementation paths.

From performance point of view our distributed solution may be seen as a tradeoff between aggregation capabilities with sliding windows and scalability. On the other hand one might argue that we lose some knowledge about specific correlations between events. In the end we agree with this argument. We do not see this property of our system to be problematic. Actually, we believe that this is property of the massive event processing itself. It is not possible to aggregate infinite amount of events with finite resources. So it seems necessary to sacrifice some aggregation capabilities to gain processing capabilities. Allowing such tradeoff on the level of theoretical CEP model gives great flexibility in the design of such applications.

6. References

- [1] Artikis, A., Etzion, O., Feldman, Z., Fournier, F. (2012). *Event processing under uncertainty*. Proceedings of the 6th ACM International Conference on Distributed Event-Based Systems - DEBS '12
- [2] Akram, S., Marazakis, M., Bilas, A. *Understanding and Improving the Cost of Scaling Distributed Event Processing Categories and Subject Descriptors*
- [3] Balis, B., Dyk, G., Bubak, M. *On-Line Grid Monitoring Based on Distributed Query Processing*
- [4] Balis, B., Slota, R., Kitowski, J., Bubak, M. *On-Line Monitoring of Service-Level Agreements in the Grid*
- [5] Barbosa, V. *An Introduction to Distributed Algorithms* The MIT Press, 1996
- [6] Biger, A., Rabinovich, Y. *Stratified implementation of event processing network*
- [7] Etzion, O., Niblett P. *Event Processing in Action* Manning Publications, 2011
- [8] Hirzel, M. *Partition and Compose: Parallel Complex Event Processing*
- [9] Huebsch, R., Hellerstein, J. M., Shenker, S. *Querying the Internet with PIER*.
- [10] Isoyama, K. *A Scalable Complex Event Processing System and Evaluations of its Performance*
- [11] Kowalewski, B., Bubak, M., Bali, B. *An Event-Based Approach to Reducing Coupling in Large-Scale Applications*.
- [12] Lee, S., Lee, Y., Kim, B., Candan, K. S., Rhee, Y., Song, J. *High-performance composite event monitoring system supporting large numbers of queries and sources*. Proceedings of the 5th ACM international conference on Distributed event-based system - DEBS '11
- [13] Luckham, D. *The Power of Events* Addison-Wesley Professional, 2002
- [14] Luckham, D. C., Frasca, B. *Complex Event Processing in Distributed Systems*
- [15] Randika, H. C., Martin, H. E., Sampath, D. M. R. R., Metihakwala, D. S., Sarveswaren, K., Wijekoon, M. *Scalable fault tolerant architecture for complex event processing systems*. 2010 International Conference on Advances in ICT for Emerging Regions (ICTer)
- [16] Renesse, R. V. A. N., Birman, K. P., Vogels, W. *Astrolabe : A Robust and Scalable Technology for Distributed System Monitoring , Management , and Data Mining*
- [17] Schilling, B., Koldehofe, B., Rothermel, K. *Distributed Heterogeneous Event Processing Enhancing Scalability and Interoperability of CEP in an Industrial Context Categories and Subject Descriptors*
- [18] Tel, G. *Introduction to Distributed Algorithms* Cambridge University Press, 2000
- [19] Vera, J., Perrochon, L., Luckham, D. C. *Chapter Event-Based Execution Architectures for Dynamic Software Systems*

- [20] Wu, E., Diao, Y., Rizvi, S. *High-performance complex event processing over streams.* Proceedings of the 2006 ACM SIGMOD international conference on Management of data - SIGMOD '06
- [21] Magid, Y., Sharon, G., Arcushin, S., Ben-Harrush, I., Rabinovich, E. *Industry experience with the IBM Active Middleware Technology (AMiT) Complex Event Processing engine* In Proceedings of the Fourth ACM International Conference on Distributed Event-Based Systems - DEBS '10

Monitorování a evaluace výukových procesů

Lucie Pekárková, Patrícia Eibenová

Fakulta informatiky, Botanická 68a, 602 00, Brno
{pekarkova, eibenova}@fi.muni.cz

Abstrakt:

Projekt MEDUSY (Multi-purpose EDUcation SYstem) se snaží využít osvědčených postupů v oblasti e-learningu a obohatit svět virtuálního vzdělávání o nové koncepty. Projekt představuje procesně orientovaný přístup k vývoji kurzů a vzdělávacích toků obecně. MEDUSY si klade za cíl být modulárním systémem (s možností připojení tradičních LMS) a poskytnout infrastrukturu pro správu a monitorování vzdělávacích procesů. V tomto článku se zabýváme teoretickými koncepty kvality a efektivity výukového procesu a zaměřujeme se na jejich měření, monitorování a evaluaci.

Abstract:

MEDUSY (Multi-purpose EDUcation SYstem) tries to leverage best practices in e-learning and enrich the world of virtual learning with new concepts. The project introduces a process-oriented approach to course development and education flow in general. MEDUSY aims to be modular and will allow plugging in traditional content-based LMS and provide infrastructure for management and monitoring of learning processes that will help to improve their quality and effectiveness. In this paper we deal with theoretical concepts of quality and effectiveness of learning processes and we focus on their measurement, monitoring and evaluation.

Klíčová slova

výukové procesy, monitorování, evaluace, vzory, Medusy

Key words

learning process, monitoring, evaluation, patterns, Medusy

1. Úvod

E-learning je jednou z atraktivních oblastí pro výzkum a vývoj. Aby se investice do této oblasti vyplatily, je nutné, abychom vytvořili robustní modely pro hodnocení e-learningu a nástroje, které budou flexibilní pro použití a zároveň konzistentní ve výsledcích. Dosavadní výzkum byl zaměřen na hledání způsobů vylepšení kvality e-learningových kurzů, jen málokdo se ale zabýval pedagogickým aspektem nebo učením. Vývoj modelů a nástrojů pro hodnocení e-learningu proto může pomoci zvýšit kvalitu výuky a ovlivnit vývoj standardů a norem.

Budoucí informační společnost bude společností vzdělávající se. Bude společností, která bude do vzdělávání investovat velké lidské a materiální prostředky. Všeobecná dostupnost informací umožní, aby se učitelé soustředili na zpracování informací, na vytváření a upevňování znalostí a rozvoj myšlení bez zbytečného memorování. Dostupnost informací vyžaduje změnu někdy vyhraněně autoritativního postoje učitele, který ob stojí před svými dobře informovanými žáky jen tehdy, bude-li jim více než dosud především partnerem a rádcem.

Dnes dochází i u nás k podstatnému rozšíření tradičních způsobů hodnocení a systematického sledování práce studentů i celých škol o mnohem efektivnější a objektivnější způsoby zjišťování a sledování účinnosti pedagogických a organizačních opatření. Nejedná se jen o technická opatření nebo o změny dílčího mechanismu řízení, ale o novou orientaci vzdělávací politiky a proměnu celého pojetí vzdělávacího systému i jeho úlohy pro rozvoj jedinců, společnosti a ekonomiky, a to ve více směrech:

- za prvé, jako projev zásadní změny pojetí pedagogických procesů od plnění předem stanovených norem všemi stejně k vytváření prostoru a podmínek pro rozvíjení individuálních talentů a nadání každého člověka (systémově inkluzivní přístup);

- za druhé, jako nezbytný důsledek decentralizace vzdělávacího systému, kdy nižší úroveň řízení (orgány samosprávy, autonomní školy) mají podstatně větší pravomoci a samy rozhodují, jak dosáhnout rámcově stanovených výsledků; mají tedy i mnohem větší vlastní odpovědnost, což vyžaduje posílení zpětné vazby evaluačními procesy;
- za třetí, jako úsilí zajistit nejen co nejvyšší kvalitu vzdělávání každého žáka, ale i efektivitu celého vzdělávacího systému.

Kromě definování různých forem hodnocení je rovněž nutné vymezit související koncepty monitorování, tj. systematického sledování vzdělávacích výsledků a vyhledávání slabých míst na různých úrovních systému. [1]

2. Evaluace vzdělávacích procesů

Existuje řada studií ukazujících výhody e-learningu. Patří sem snížení nákladů na výuku, schopnost zasáhnout významně vyšší počet studentů, zpracovat rozsáhlejší množství vědomostí, efektivnější řízení vzdělávacích procesů, v komerční sféře pak zvýšení spokojenosti zaměstnanců či snížení jejich fluktuace. V dnešní ekonomice nám však nestačí pouze důvěřovat těmto studiím a tvrzením, každá investice musí být vyhodnocena a její přínosy musí být očividné. Z toho důvodu řada společností neinvestovala do vzdělávání, poněvadž byly obtížně měřitelné výhody tohoto investování. Postupně však vzniká metodologie, jak měřit efektivitu e-learningu a jak aplikovat výpočet návratnosti investic i v e-learningu.

2.1 Vymezení pojmů

Pojem evaluace, hodnocení, kvalita i pojem efektivita jsou v pedagogice používány v různých významech. Často dochází k překrývání významů obou pojmů, a proto vždy záleží na kontextu, ve kterém se používají.

2.1.1 Evaluace, hodnocení, evaluace vzdělávání

Pojem **evaluace** je v tomto kontextu užit jako zastřešující termín, který označuje systematické shromažďování, třídění, posuzování, vyhodnocování a analýzu prvků vzdělávacího procesu s cílem zvýšit jejich kvalitu i efektivitu. Zdůrazňuje systematickost, strukturovanost a plánovitost procesu, zejména jeho propojenost s rozvojem a zlepšováním výukových procesů, slouží jako zpětná vazba pro všechny účastníky a vztahuje se na všechny fáze vzdělávacího procesu.

Evaluace je významnou, závěrečnou součástí vzdělávacího procesu. Jedná se o pokus získat informace (zpětnou vazbu) o účincích vzdělávacího procesu a ocenit hodnotu tohoto vzdělávání ve světle získané informace. Zpětná vazba je součástí vyhodnocovacího procesu nebo také neustálým hodnocením potenciálu a produktivity všech pracovníků a efektivnosti prostředků vynaložených na vzdělávání.

Pojem **hodnocení** není chápán jako zaměnitelné synonymum, ale je užíván jednak v užším smyslu pro hodnocení jednotlivých studentů příp. i učitelů (obdobu rozdílu anglického *evaluation* a *assessment*), jednak v širším smyslu pro volnější vyjádření vlastního aktu/procesu hodnocení v běžné školní praxi.

Aby splnil svůj úkol zajišťovat kvalitu i efektivitu, je mechanismus evaluace vzdělávání vzdělávacích výsledků i práce škol nezbytně komplexní:

- zasahuje všechny úrovně vzdělávacího systému (jednotlivého žáka, celé školy, vymezené skupiny škol - jednoho druhu škol, jednoho regionu - i celého státu), umožňuje agregaci na jeho každé úrovni i porovnání na mezinárodní úrovni, mezi celými vzdělávacími systémy;
- využívá a propojuje různé, vzájemně se doplňující formy evaluace – interní i externí, subjektivní i objektivně měřitelné, sumativní i formativní, umožňuje porovnání mezi různými subjekty vzhledem k jejich individuálním a specifickým podmínkám a možnostem, a umožňuje i porovnání v čase, tedy sledování vývoje. [1]

2.1.2 Kvalita a efektivita vzdělávacích procesů

Pojem **kvalita** se ve vzdělávání vyskytuje nejčastěji pro vyjádření stavu, který je optimální, žádoucí, ideální, tedy a priori pozitivní. Kvalitou (výuky, vzdělávacích procesů, vzdělávacích institucí, vzdělávací soustavy) se pak rozumí „žádoucí (optimální) úroveň fungování a/nebo produkce těchto procesů či institucí, která může být předepsána určitými požadavky (např. vzdělávacími standardy), a může být tudíž objektivně měřena a hodnocena“ [2]. Obecně je tedy kvalita vyjádřením nějakého stavu.

Z diskuzí o nejlepších strategiích pro kvalitní e-learningové kurzy vychází najevo, že musí být zaměřené na studenta. To zahrnuje i nutnost jasně vyjádřit studentovy potřeby před zahájením kurzu. Důležitými aspekty se tak stávají uvědomění si učebního stylu, individuálních učebních preferencí a sociálních potřeb každého studenta.

Je důležité poznamenat, že kvalita výukového procesu není něco, co by student dostával od učitele (jakožto poskytovatele e-learningového obsahu), ale spíše je tvořena procesem spolupráce mezi studentem a učebním prostředím. To znamená, že kvalita výukového procesu není pouze výsledkem produkčního procesu výukové instituce. Kvalita spíše motivuje, umožňuje studovat a musí být studentem definována až v konečné fázi výukového procesu. [3]

Efektivita (vzdělávání, vzdělávacích procesů a vzdělávacích institucí) zahrnuje dva významy – jejich účelnost (tj. vhodnost užití prostředků pro dosažení stanovených cílů za daných podmínek) i jejich účinnost (danou poměrem dosažených výsledků k vynaloženým zdrojům – personálním, materiálním i finančním). Efektivita odkazuje jednak na účinky, výsledky, následky či důsledky, a také na jejich zdroj, původ, příčiny. Efektivita je tedy obecně vyjádřením určitého vztahu (často mezi výsledkem a tím, co tento výsledek způsobilo, popřípadě ovlivnilo). Ve výuce se jedná o konečné výsledky jako například znalosti žáka na konci školního roku, počty přijatých studentů na vysokou školu. [4]

2.2 Metody měření efektivity výukových procesů

Základem měření efektivity, evaluace výukového procesu, je určení, jaké informace a jaké množství informací je nutné shromáždit, a také výběr správné metody, která by zajistila validitu zjištěných informací. Hodnocení mohou provádět účastníci (hodnotí kvalitu výukového procesu), organizátor výukového procesu (hodnotí úroveň výkonu učitele a průběh procesu z hlediska získaných zkušeností), vedení organizace (hodnotí přínos výukových procesů pro řešení problému), učitel (hodnotí vlastní výkon z hlediska dosažených výsledků u účastníků).

Při hodnocení výukového procesu dominují průzkumy spokojenosti. Tento výzkum je daleko snadněji měřitelný než výzkum zaměřený na výsledky. Pro co největší měřitelný efekt vzdělávací aktivity je důležité balancovat mezi hodnocením spokojenosti, tedy něčím, co má iracionální podtext, a zároveň tím, co je faktickým, ale zároveň ve vztahu ke vzdělávací aktivitě. Obvykle platí, že metody s vysokou reliabilitou nemívají příliš velkou validitu. Abychom se co nejvíce vyvarovali nedokonalostem měřících nástrojů, je dobré je navzájem kombinovat, aby se navzájem doplňovaly.

2.2.1 Kirkpatrickův model

V literatuře [5] se můžeme setkat s široce akceptovanou metodou měření efektivity školicích programů, vyvinutou Donaldem L. Kirkpatrickem již v roce 1959. Kirkpatrickův model zahrnuje 4 stupně vyhodnocení:

- Stupeň 1: Reakce – Jak studenti reagují na školení?
- Stupeň 2: Výuka – Kolik se toho naučili?
- Stupeň 3: Chování – Jak se změnilo jejich chování?
- Stupeň 4: Výsledky – Jaký efekt mělo školení pro organizaci?

K těmto 4 stupňům byl přidán pátý stupeň:

- Stupeň 5: Návratnost investic – Převážily výsledky ze školení jeho cenu?

Pro měření efektivity výuky z pedagogického hlediska, což v podstatě odpovídá dle výše uvedeného členění stupni 2, je důležité, aby jakýkoliv konkrétní systém výuky fungoval v relativně ustálených podmínkách. Vnesení jakékoliv změny do tohoto systému má za následek způsobení odchylky v podobě zvýšení či snížení efektivity. Změnou může být chápáno užití jiné organizační formy, výukové metody či libovolného materiálního didaktického prostředku. Při jiném úhlu pohledu se celková efektivnost výuky odvíjí již od její přípravy, vlastní realizace a závěru. Nelze jednoznačně určit, která z těchto fází je pro výslednou efektivnost výuky důležitější a proto je nutné všem věnovat náležitou pozornost.

Evaluace je jen poslední částí vzdělávacích procesů, který obsahuje další důležité aktivity jako např. analýzu vzdělávacích potřeb. Kirkpatrick proto zdůrazňuje, že úspěch jednotlivého kurzu tkví v systematické práci na všech dílčích částech procesu vzdělávání a souvisí s pečlivě formulovanými cíli a přípravou. Měření efektivity je pak smysluplné, pokud z něj získané informace zpětně ovlivňují celou přípravnou část procesu. Na měření efektivity proto můžeme pohlížet jako na pyramidu, na jejímž vrcholku je čtvrtý stupeň měření (výsledky) a na proces vzdělávání jako na kruhouvou cestu od získávání informací o potřebách vzdělávání, jejich využití při přípravě školení, realizaci vzdělávací akce a její vyhodnocení, které opět přináší informace o potřebách vzdělávání a zpětnou vazbu o efektivitě celého procesu.

2.2.2 Exaktní měření efektivity výukového procesu

V podstatě jsou známy následující možnosti jak exaktně měřit efektivitu výuky. V první řadě lze posuzovat to, co si studenti měli osvojit a co si skutečně osvojili. Srovnáváme tak referenční skupinu, za předpokladu, že cíle výuky jsou v obou případech stejné a odhlížíme od rozdílného složení studentů i podmínek studia. Při tomto posuzování postačuje provedení jen jednoho měření. Nejvhodněji k tomu poslouží didaktický test. Výpočet je možné provést podle vztahu [6]:

$$E_I = \frac{V_{post}}{V_{max}} * 100 \quad (\text{pro jednotlivce})$$

$$E_G = \frac{\sum^N \left(\frac{V_{post}}{V_{max}} * 100 \right)}{N} \quad (\text{pro skupinu})$$

kde V_{post} je výsledek, který byl naměřen testem po proběhnutí výuky, V_{max} je nejvyšší možný dosažitelný výsledek. Nevýhodou tohoto měření ovšem je, že není rozlišováno to, co studenti věděli ještě před sledovaným edukačním procesem, od toho, co se skutečně ve výuce naučili (chybí porovnávání výchozího a konečného stavu téže skupiny studujících). K odstranění tohoto nedostatku je možné použít vztah:

$$E_I = \frac{V_{post} - V_{pre}}{V_{max}} * 100 \quad (\text{pro jednotlivce})$$

$$E_G = \frac{\sum^N \left(\frac{V_{post} - V_{pre}}{V_{max}} * 100 \right)}{N} \quad (\text{pro skupinu})$$

Nevýhodou uvedeného postupu je ovšem nutnost provést měření pomocí dvou testů. To lze odstranit užitím statistických metod, např. jednofaktorové analýzy rozptylu.[7]

2.2.3 Evaluace vzdělávacích procesů

Pokud chceme vytvořit hodnotící proces, který bude efektivní, měli bychom mít na paměti, tyto charakteristiky efektivního hodnotícího procesu [8]:

- flexibilní, vhodný pro různé formy výuky: efektivní hodnotící proces by měl umět vyhodnotit širokou škálu výukových kurzů. Měl by hodnotit výukový kurz z různých úhlů pohledu, aby bylo možné hodnotit jak technické, tak soft-skills kurzy, jak formální, tak neformální výuku;
- jednoduchý, jednoduše implementovatelný: dobrý proces je tak jednoduchý, že mu každý zapojený účastník rozumí, manažeři věří výsledkům a účastníci jej zvládnou;

- spolehlivý, schopný přesně měřit nebo předvídat výsledky: na konci evaluace je nutné mít smysluplné výsledky, které jsou konzistentní nebo jednoduše vysvětlí rozdíl v efektivitě výukového procesu;
- ekonomický, nezahrnující vysoké výdaje: dobrý hodnotící proces by neměl být finančně, časově a personálně náročný.

3. Monitorování vzdělávacích procesů

Na hodnocení a monitorování vzdělávacích procesů je nutno nahlížet jako na celek, na propojené aktivity. Data získaná z monitorovacího procesu slouží jako podklad pro evaluaci a optimalizaci procesů, proto je důležité znát procesy, určit monitorovací indikátory, a mít tak možnost reakce na běžící proces.

3.1 Motivace monitorování vzdělávacích procesů

Monitorování vzdělávacích procesů provádíme z různých důvodů. Potřebujeme získat data, která jsme si na začátku definovali a pomocí nich určit, zda byl proces efektivní a jaká je jeho kvalita. Tato data nemusíme vyhodnotit jenom po skončení procesu, ale můžeme s nimi operovat i v době, kdy běh procesu ještě nebyl ukončený, například za účelem okamžité reflexe. Vzdělávací procesy mají specifický průběh, vzhledem k jejich specifickým účastníkům. Proto také data z každého spuštěného procesu budou různá, a abychom je mohli vyhodnotit korektně, musíme přihlížet i ke všem doprovodným okolnostem.

Dalším důvodem pro monitorování výukových procesů je vyhodnocení ukončeného procesu a sběr dat pro následnou optimalizaci, odstranění chyb a slabých míst, případně zvýšení efektivity v další iteraci. Jiným argumentem pro monitorování je sledování procesů, abychom jednoduše určili jejich stav (což pomáhá při identifikaci problému a usnadňuje reakci), určili výkon a efektivitu běžících procesů a také hromadně sbírali data pro vytváření statistik.

Typickým příkladem monitorování a vyhodnocování bývá statistika známek po ukončení kurzu. Vyučující z tohoto hodnotícího kritéria vyvodí důsledky a promítne je do další iterace kurzu. Motivací pro komplexnější monitorování je poptávka po detailnějších datech, jejichž znalost pomůže nahlížet na proces z jiné perspektivy, proces upravit, přizpůsobit množině účastníků a přinést lepší výsledky. Za předpokladu, že je výsledné hodnocení součtem více faktorů, vznikajících během spuštěného procesu, můžeme sledováním procesu jednotlivé části pojmenovat a zjistit, proč a jak ovlivnily výsledek.

Monitorování předchází identifikace samotných procesů. Znalost těchto procesů je stěžejní pro definování monitorovacích ukazatelů a faktorů. Není možné předpokládat, že člověk ve funkci lektora bude mít detailní povědomí o běhu procesu a určení monitorovacích indikátorů. Propojení samotných procesů s monitorovacími indikátory na základě množiny podobných faktorů je základ pro automatizaci nejen monitorování procesů, ale také jejich samotného běhu.

3.2 Klíčové indikátory výkonnosti - Key Performance Indicators

Klíčové indikátory výkonnosti - Key Performance Indicators (dále KPI) slouží pro kvantifikaci výstupů procesů a stanovení úrovně efektivity procesu. KPI se stanovují předem, na konci procesu pak probíhá zhodnocení, jestli se podařilo dosáhnout stanovených cílů. Kritické přitom je, že ještě před spuštěním procesu se ví co je důležitým prvkem a na co se má zaměřit pozornost, co pomáhá bránit odklonění od původních cílů.

KPI jsou vyjádřeny čísly, jsou tedy jednoduše interpretovatelné a nezaměnitelné. Podstatné vlastnosti KPI jsou účelnost, jednoznačnost, zjištělnost a interpretace. Na tyto vlastnosti se klade důraz již při návrhu KPI a stávají se tak jakýmsi standardem pro zhmotnění definic a cílů. Jakmile proces běží, definované indikátory podléhají měření s možností průběžné kontroly a reflexe. Po skončení procesu jsou vyhodnoceny porovnáním s definovanými cíli v etapě návrhu procesu.

3.3 Monitorování v LMS systémech

Monitorování procesů v LMS (Learning Management System) systémech je v dnešní době zaběhlým jevem. Vyučující mají možnost sledovat, jak studenti prochází procesem, jak se jim daří plnit zadané úkoly, atd. I tady existují KPI, v zaběhlém vnímání již absolutně automatické (např. hodnocení, známka). Před začátkem kurzu (spuštěním procesu) se stanoví bodové hranice, po skončení (procesu) se pak vyhodnotí dosažené body, vytvoří se statistiky známek, určí se průměry a obtížnosti. Vyhodnocování může probíhat i průběžně, což slouží k přizpůsobení kurzu konkrétní skupině studentů.

Hodnocení není jediným kritériem pro určení úspěšnosti procesu (kurzu). Kritéria jsou komplexní množiny dat, přičemž platí čím větší detail, tím přesnější výsledek. Jak bylo napovězeno u případu průběžného vyhodnocování, někdy je třeba přihlídnout k průběžným výsledkům monitorování a přizpůsobit průběh procesu. Výsledná známka je pouze výstupem mnoha faktorů, a pokud je nezachytíme přímo při běhu, po ukončení procesu budeme mít pouze výsledek s neznámými proměnnými. Také naměřené výsledky s velkým rozsahem, ale bez souvislostí, nabídnou pouze anonymní obecná data nepoužitelná v konkrétní situaci.

3.4 Vzory ve vzdělávacích procesech a vztah k monitorování

Vzory ve vzdělávacích procesech jsou oblastí výzkumu již několik let. Jejich hledáním a využitím se v devadesátých letech minulého století zabýval Pedagogical Patterns Project¹⁵, který definuje vzor jako výsledek zobecnění nejlepších zkušeností v určité oblasti, přičemž je podstatné problém neizolovat a sledovat jej v konkrétním kontextu. Projekt měl pomoci učitelům vytvářet výuku a přinést tak lepší výsledky ve vzdělávání studentů.

Ještě blíže k námi zkoumané oblasti se nachází The Person-Centered e-Learning Pattern Repository¹⁶, který úspěšně propojil přístup zaměřený na člověka s tak obecnou disciplínou jakou je definování vzorů.

Budoucnost edukačních procesů by mohla být právě v zakomponování základů přístupu zaměřeného na člověka do vytváření vzdělávacích vzorů a definování kritérií hodnocení procesů. Tímto přístupem se z výsledků monitorování získají více než jenom prázdná anonymní data. Čím více obecný výsledek je, tím menší výpovědní hodnotu obsahuje.

3.4.1 Použití vzorů při monitorování procesů

Předpokladem využití vzorů při monitorování procesů je provázání vzoru se stanovenými KPI. Každý vzor obsahuje množinu KPI, které jsou jeho pevnou součástí a jasně definují, co se ve vzoru bude měřit. Uživatel, resp. učitel, který se s procesem seznámí, bude mít ještě před jeho spuštěním informace o definovaných kritériích měření a možnostech, které může vyvozovat z měření. Data, vyplývající přímo z monitorování, by měla být jednoduše identifikovatelná a ihned použitelná, jelikož mohou sloužit k okamžité reflexi.

Případné nedostatky v definovaných KPI vyplynou po skončení procesu a budou sloužit pro následnou optimalizaci procesu. V další iteraci životního cyklu vzdělávacího procesu bude problém již odstraněn. Řešení může být zakomponováno do vzoru jako další možnost výběru nebo bude sloužit v konkrétní specifikaci této instance procesu.

Úlohou učitele by také nemělo být určování, co znamená kvalitní vzdělávací proces a co nikoliv. Zodpovězení těchto otázek by mělo spočívat v dokonalé definici vzdělávacích vzorů, včetně určení KPI.

4. Nasazení monitorování a evaluace výukových procesů

¹⁵ <http://www.pedagogicalpatterns.org/>

¹⁶ <http://elearn.pri.univie.ac.at/pca/?show=projects&n=Patterns>

4.1 Projekt MEDUSY

Projekt Medusy (Multi-purpose EDUcation SYstem¹⁷) vznikl v rámci průmyslového partnerství mezi laboratoří Lasaris na Fakultě informatiky a společností Red Hat Czech Republic. Projekt se snaží využít některých z osvědčených postupů v oblasti e-learningu, implementovaných v moderních LMS systémech a obohatit svět e-learningu o nové koncepty. Projekt představuje procesně orientovaný přístup k vývoji kurzů a vzdělávacích toků obecně. MEDUSY si klade za cíl být modulárním systémem s možností připojení tradičních LMS zaměřených na správu obsahu a poskytnout infrastrukturu pro správu a monitorování vzdělávacích procesů.

Cílem projektu MEDUSY je poskytnout vhodné prostředí pro celý životní cyklus procesu, tedy pro modelování, spouštění, monitorování a optimalizaci. V rámci projektu vzniklo již několik diplomových prací, na jejichž základech bude stavět další výzkum.

4.2 Dosavadní výsledky MEDUSY a jejich propojení s monitorovacím nástrojem

Projekt MEDUSY je postaven na platformě Activiti BPM suite. Activiti však v nynější verzi neobsahuje komponentu pro monitorování, nejvhodnějším řešením bude proto vytvoření vlastního monitorovacího nástroje pro projekt. Nástroj ve formě komponenty bude kooperovat s dosavadními nástroji vytvořenými v projektu a bude přizpůsoben přímo na míru potřeb projektu. Vytvořené vzory společně s požadavky na monitorování totiž vyžadují velmi specifické prostředí.

Monitorovací nástroj bude napojený na již existující engine vytvořený v rámci diplomové práce Mgr. Daniela Továrnáka¹⁸. Monitorování bude probíhat v samostatné komponentě, ale bude sledovat procesy spuštěné v tomto procesním engine. Kooperace těchto dvou komponent bude proto klíčová.

Současně s implementací komponenty na monitorování procesů bude vytvořen prototypu vzoru vzdělávacího procesu i s definovanými KPI. Je možné využít již existující vzory ze základního repozitáře vytvořeného Mgr. Jiřím Novákem. Ten ve své diplomové práci¹⁹ definoval základní vzory, které se po úpravě a přidání KPI mohou použít jako testovací data pro monitorování a nakonec postaví základ pro repozitář vzorů projektu.

5. Závěr

Výukové procesy představují složité, otevřené systémy, na které působí vlivy vnějšího prostředí, a proto mohou fungovat s různou efektivností. Jednou z možností jak ji měřit, je využití didaktických testů ve spojení se statistickými výzkumnými metodami. Hlavními kritérii efektivnosti tak mohou být čas, vynaložená energie (studenta i učitele) nebo výsledky výukových procesů.

Vstupní data k monitorování se liší s každou skupinou studentů, je proto klíčové při definování KPI v rámci vzorů pokrýt všechny možnosti. Systém, který je efektivní pro jednu skupinu nebo i jednotlivce, nemusí být stejně účinný pro jinou. Ovšem naopak to automaticky neznamená, že pokud je proces neefektivní, je třeba ho upravit. Navrhovaný monitorovací systém by měl pokrýt i tyto aspekty a nahlížet na data v potřebném kontextu.

Co se týče možností různých LMS v oblasti statistik a analytických reportů, není tato oblast zatím velmi podporována. Některé LMS nabízí možnost jednotlivě zjistit, kdo nebo kdy přistupoval ke studijním materiálům, kdy byl naposledy přihlášen nebo kdy odevzdal úkol. To ale pro analýzu chování studentů, zpětnou vazbu a vůbec analytický report celého kurzu nestačí.

Je tedy nutné vytvořit sadu KPI pro výukové procesy, resp. vzory, které budou obsahovat jak tvrdá data (např. známky, čas strávený učením), tak data měkká (např. spokojenost studenta s kurzem, lektorem) a ty pomocí nového modulu v projektu MEDUSY monitorovat, vyhodnocovat a dále s nimi pracovat.

¹⁷ http://sourceforge.net/apps/mediawiki/medusy/index.php?title=Main_Page

¹⁸ http://is.muni.cz/th/172673/fi_m/thesis.pdf

¹⁹ http://is.muni.cz/th/172704/fi_m/dp.pdf

6. Literatura

- [1] *Rámcový projekt monitorování a hodnocení vzdělávání*. [online]. 2003 [cit. 2012-10-30]. Dostupné z WWW: <<http://aplikace.msmt.cz/PDF/JORamcovyprojektmonitorovaniahodnocenivzdelavani.pdf>>.
- [2] Průcha, J. *Pedagogická evaluace*. Brno : MU CDVU, 1996, s. 27.
- [3] Ehlers, U-D. Quality in e-learning from a learner's perspective. *European Journal of Open and Distance Learning*, 2004-I, 2004. Dostupné z WWW: <<http://www.eurodl.org/index.php?tag=120&article=230&article=101#1>>.
- [4] Janíková M., Vlčková K. *Výzkum výuky: tématické oblasti, výzkumné přístupy a metody*. 1. vyd. Brno: Paido, 2009, s. 63-82. Pedagogický výzkum v teorii a praxi, sv. 13. ISBN 978-80-7315-180-5.
- [5] Petlák, E. *Všeobecná didaktika*. 2. vyd. Bratislava: IRIS, 2004. 311 s. ISBN 80-89018-64-5.
- [6] Poláková, E., Štefančíková, A. Meranie efektívnosti dištančného štúdia. In *Technika – informatyka, edukacja*. Rzeszów: Uniw. Rzeszowski, 2005. s. 126–133. ISBN 83-88845-55-1., str. 130.
- [7] Dostál, J. Pedagogická efektivita off-line learningu v celoživotním vzdělávání. In: *Klady a zápory e-learningu na menších vysokých školách, ale nejen na nich: Konference : Praha, 23. května 2008*. Vyd. 1. Praha: Soukromá vysoká škola ekonomických studií, 2008, s. 56-64. ISBN 978-80-86744-76-6. Dostupné z: <http://www.svses.cz/projekty/konference/e_learn/sbornik_%203153.pdf>.
- [8] Horton, W. K. *Evaluating E-learning*. Vyd. 3. Alexandria USA: ASTD, 2004, 125 s. ISBN 15-628-6300-2.

Dohledové systémy

Tomáš Pitner

Masaryk University, Faculty of Informatics, Lab Software Architectures and IS
Botanická 68a, 602 00 Brno, Czech Republic
tomp@fi.muni.cz

Abstrakt.

Monitorovací a dohledové systémy jsou podstatnou součástí IT infrastruktury podniků a dalších organizací.

Abstract.

Monitoring and surveillance systems are an essential part of the IT infrastructure, companies and other organizations.

Klíčová slova

monitoring, dohledové systémy

Keywords

monitoring, surveillance systems

1. Úvod

Dohledové a monitorovací systémy jsou v dnešní době hojně nasazovány na sledování funkčnosti, spolehlivosti a výkonu v nejrůznějších počítačových a jiných infrastrukturách, jako jsou např. budovy a výrobní kapacity.

Dohledový a monitorovací systém je charakteristický tím, že průběžně:

- shromažďuje data ze sledovaného prostředí,
- vyhodnocuje tato data a identifikuje události
- upozorňuje uživatele na důležité události,
- nebo je schopen reagovat automaticky.

Systémy se používají se ke sledování provozu objektů v nejširším smyslu, od jedné počítačové aplikace, přes jednotlivé počítače a další zařízení až po velké infrastruktury (budovy, sítě). Předmětem výzkumu a vývoje je také integrace dílčích dohledových systémů a propojení s informačními systémy, např. geografickými nebo ERP systémy. Díky nasazení dohledového systému máme možnost:

- rozpoznat neobvyklé chování,
- zjistit, zda nedošlo k porušení pravidel chování,
- detekovat odlišnosti od dlouhodobého, normálního profil,
- identifikovat odlišnosti od podobných objektů (např. sousední místnost nebude normálně mít příliš odlišnou teplotu),
- měřit indikátory výkonu v reálném čase vč. dodržování SLA,
- objevovat fakta, která nejsou zjevná nebo
- zkrátit čas detekce a
- snížit čas potřebný na odhalení původní příčiny chybného chování,
- zabránit regresi chyb při zavádění změn,
- snížit náklady na zaměstnance – vyžadují méně kvalifikovaný personál.

2. Aplikační oblasti dohledových systémů

2.1 Správa budov

První oblastí zájmu skupiny dohledových systémů v laboratoři Lasaris je problematika správy budov (facility management). Studenti laboratoře jsou angažováni v reálném provozu na Oddělení systémů inteligentních budov Ústavu výpočetní techniky MU a rovněž Univerzitního kampusu Bohunice (UKB). Obecně jde o dohled a analýzy nad procesy správy budov a jednotlivými technologiemi, jimiž je budova vybavena.

Masarykova univerzita s devíti fakultami, téměř 44 tis studenty a 4500 zaměstnanci je druhou největší v České republice. Disponuje více než 250 budovami s více než 20 tis místnostmi o celkové ploše více než 350 tis m². Kampus Masarykovy univerzity je největším VŠ areálem v ČR, s více než 100 tis M2 plochy a 30 budovami (přibývají další). Budovy jsou plně vybavené provozními technologiemi jako topení, chlazení, klimatizace, odvod vody. K dohledu patří zabezpečení: požární hlásiče, řízení přístupu, kamery. Do technologií zde patří i audiovizuální technika, osvětlení, napájení a nakládání s odpady. Monitorovací systém produkuje denně řádově 100 tis. záznamů.

V rámci diplomové a nyní i dizertační práce Adamy Kučery byly identifikovány hlavní aplikační oblasti vzory použití metod zpracování komplexních událostí v řízení budov.

2.2 Počítačová bezpečnost sítí

Pro zabezpečení rozsáhlé počítačové sítě, jakou disponuje např. univerzita, lze použít různé přístupy, v zásadě se pohybující mezi dvěma póly: restriktivním a liberálním:

- Více omezující je přístup s nasazením přísně nastavených firewallů, v akademickém prostředí obtížně prosaditelné až nepoužitelné.
- Liberální, méně restriktivní formou je monitoring a následná reakce v případě porušení pravidel uživatelem zevnitř nebo útoku zvenčí.

Pro implementaci liberálního přístupu se na Masarykově univerzitě využívá metod založených na sledování toku v síti (netflow-based). Směr, který se před lety zdál jako nerealistický, se postupně se vzrůstajícími technickými a výpočetními možnostmi stává reálnou alternativou, kterou volí firmy a další větší instituce k zabezpečení své sítě.

Výhodou je minimální obtěžování řádně se chovajícího uživatele, odpadá nutnost příliš restriktivně nastavovat systémové politiky či instalovat speciální klientský software na koncová zařízení. To jsou důvody, proč tým CSIRT (Computer Security Incident Response Team) Masarykovy univerzity nasadil tento přístup v univerzitní síti. Výsledkem byl např. úspěšně odhalený botnet Chuck Norris [XXX].

Systémy založené na sledování toku a detekci mimořádných událostí (např. útoků zvenjšku) jsou zatím nasazovány sice hierarchicky na příslušných uzlech univerzitní sítě, ale do budoucna je bude třeba propojit i se sítěmi jiných provozovatelů, v první fázi v rámci akademické sítě CESNET. Díky tomu bude možné informace o zachycených incidentech sdílet a tak pomoci ochránit další síť. Rovněž je patrné, že bez výměny informací o podezřelém chování nebude možné některé z útoků vůbec rozpoznat, resp. odhalit včas.

2.3 Rozvodné sítě

Smart-gridy, chytré produkční a rozvodné sítě elektřiny jsou možným příspěvkem k zajištění energetické bezpečnosti v budoucnu. Velké výrobní a distribuční společnosti experimentují v posledních letech s jejich provozem. Typickým prvním krokem k proměně běžné distribuční sítě na „inteligentní“ je vybavení odběrních míst „chytrými elektroměry“, *smart-metery*. Jelikož v národních rozměrech je počet odběrních míst vysoký, v řádech milionů, je třeba do této velikosti škálovat i navržené řešení na sběr a zpracování dat z smart-meterů jakožto koncových měřidel. Objem sbíraných dat bude v řádech 10 TB ročně, což není velké na ukládání, ale může při nevhodném sestavení architektury být náročné na výpočetní výkon.

Škálovatelnost stávajících řešení sběru a zpracování dat, odzkoušených na řádově desítkách tisíc smart-meterů, je předmětem experimentů, na nichž se studenti Lasaris rovněž podílejí. Předmětem

vlastního využití dat ze smart-meterů je kromě měření spotřeby zejména detekce neoprávněné manipulace se zařízením, detekce výpadků dodávek, případně sledování odběrových křivek.

2.4 Simulace dohledových systémů na cloudech

Díky projektu CERIT-Scientific Cloud a mnohaletou tradicí v distribuovaném počítání a síťových technologiích disponuje Masarykova univerzita datovými úložišti a výpočetní silou nabízenou pro experimentální využití. Cílovým stavem v 2013 by mělo být 3500 procesorových jader a až 3,5 PB úložného prostoru dostupné prostřednictvím Národní gridové infrastruktury.

Z hlediska průníků s oblastí dohledových systémů nacházíme možnosti jak v oblasti sledování samotné virtualizované poskytované infrastruktury (výše zmíněný projekt Heimdal), tak v jejím využití pro simulace sledované jiné infrastruktury. Simulační experimenty v prostředí cloudu přinášejí očividné výhody:

- Lze provádět i experimenty ve vnějším prostředí příliš rizikové, např. „pěstovat si“ na virtuálních strojích v izolované virtuální síti počítačové viry.
- Prostředí pro experiment lze podstatně snadněji připravit a nakonfigurovat.
- Experiment může být pozastaven, jeho stav zmražen a následně znovu spuštěn. Díky uložení obrazů jednotlivých virtuálních strojů je možné pokračovat různými cestami a vracet se k předchozím „zapamatovaným“ stavům.
- Prostředí se snadněji škáluje než ve světě fyzických počítačů.

2.5 Virtualizované výpočetní systémy

Připravovaná dizertace D. Tovarňáka se věnuje problematice monitoringu především virtualizovaných výpočetních infrastruktur, tj. cloudů. Tato oblast je charakteristická hlavně:

- zájmy více stran na monitoringu (provozovatel cloudu/IaaS, platformy/PaaS, konkrétní aplikace a konečně i uživatel této aplikace), z toho plynoucí různé požadavky na to, co sledovat (multi-tenancy);
- nutností poskytovat informace striktně izolovaně pro každého příjemce (isolation);
- potřeba držet co nejnižší režii monitorovacích procesů ve smyslu spotřeby strojového času a paměti (low overhead).

Blíže k navrženým konceptům a popisu systému Heimdal v [XXX].

2.6 Odhalování podvodů

Studenti laboratoře se spolu s průmyslovým partnerem MycroftMind, a.s. podíleli na budování systému pro detekci podvodů na síti čerpacích stanic. Typickým scénářem podvodníka je prodej jiného než autorizovaného paliva dodaného majitelem sítě nebo čerpání (řidšího) paliva poté, co se cisterna zahřeje na slunci. Úloha odhalování podvodů je specifická tím, že ji lze těžko řešit bez podpory komplexního zpracování událostí, protože při identifikaci možného podvodu je třeba sledovat události z různých dílčích systémů, které navíc proudí s různým časovým odstupem. Zatímco s malou prodlevou lze zjistit pohyb hladiny paliva v nádržích čerpací stanice nebo jednotlivá tankování do vozidel ze stojanů, tak fakturace dodaného paliva probíhá s řádově denním zpožděním. Očekávané úspory pro středně velký řetězec mohou představovat i desítky milionů korun ročně.

Bylo navrženo řešení zahrnující implementace vybraných návrhových vzorů pro detekci jednotlivých komplexních událostí. Systém byl prototypově implementován.

2.7 Průmyslová výroba

Doktorandi laboratoře Lasaris rovněž experimentovali s rozšířením výrobního informačního systému PHARIS od spol. UNIS pro potřeby rozšířeného monitoringu výroby ve strojírenských podnicích.

Monitorovací systém sleduje jednotlivé stroje z hlediska počtu aktivních (pracovních) cyklů, eviduje, kolik cyklů bylo provedeno. Sleduje, kdo z operátorů je kdy přihlášen a přiřazuje události k nim. Vidí,

jaká operace momentálně na stroji probíhá, tzn. zda se jedná o běžnou práci, odstávku z důvodu údržby nebo poruchy. Jelikož dohledový systém je schopen sledovat tok událostí, může odvodit profil daného stroje, tzn., zda nemá příliš prostojů z důvodu poruch, zda řešení jednotlivé poruchy netrvá příliš dlouho apod.

3. Výzkumná témata v oblasti dohledů a zpracování událostí

Oblast dohledů a zpracování komplexních událostí je perspektivní jak pro praktické aplikace, tak pro řešení úkoly více výzkumného charakteru. Nastíníme si nyní ty podstatné.

3.1 Obtížná udržitelnost aplikací CEP

Vývoj aplikací CEP není jednorázovým procesem, ale probíhá po celou dobu životního cyklu. Přidaná hodnota řešení na bázi CEP je právě v tom, že aplikace „roste“ se svými uživateli inkrementálním doplňováním pravidel na odhalování nových událostí, jejich agregování apod. tzv. CEPyramida se zvyšuje. Proces přidávání nových pravidel, jakož i integrace dalších zdrojů dat/událostí vyžaduje údržbu mnoha verzí programů CEP, a podobně i architektury a konfigurací celých řešení. Dosavadní prostředky na řízení vývoje SW produktů nelze v dostatečné míře zapojit do podpory těch fází životního cyklu, kde se na vývoji podílí koncoví uživatelé, obvykle stále v součinnosti s vývojáři a doménovými experty. Varianty tohoto problému zahrnují rovněž situace, kdy je CEP řešení poskytováno jako služba, tj. často ve více verzích týmh nebo různým klientům souběžně.

3.2 Sdílení znalostí a vzorů v aplikacích CEP

V současnosti je tvůrcům aplikací CEP k dispozici několik alternativních nástrojů či platform, jak s uzavřeným, tak s otevřeným zdrojovým kódem. Prakticky každá používá vlastní jazyk pro zachycení detekčních vzorů/pravidel, jejich výpočetní modely disponují jiným repertoárem datových typů, jsou různě efektivní a škálovatelné. Situace je obdobná jako před lety u modelování informačních systémů, dokud neexistoval jazyk UML. Pro komunikaci mezi koncovým uživatelem, doménovým znalcem a vývojářem či poskytovatelem služby chybí jednotící, vizuální jazyk pro sdílení znalostí a navržených vzorů jak mezi jednotlivými lidskými aktéry, tak případně mezi systémy CEP.

3.3 Detekce příčin chyb a jejich následky

Dohledové systémy by měly být nejen schopny detekovat, že určitá funkcionalita sledovaného systému vykazuje závadu (nefunguje) nebo systém nesplňuje jistý mimofunkční požadavek (dostupnost, odezva), ale měly by rovněž napomoci odstranit příčinu této chyby. Nezbytným předpokladem je lokalizace prapůvodní příčiny chyby („root cause“), případně souběhu více chyb. To zahrnuje možnost „post-mortem“ analýzy (tj. historie) předchozího dění ve sledovaném systému.

3.4 (Polo)-automatická reakce na chyby

Jestliže sledovaný systém vykazuje chyby, které lze „odstranit“ standardní posloupností (manuálních) kroků, může automatizace těchto kroků přinést značné úspory. Takovéto řešení však s sebou přináší některá úskalí. Namátkou lze zmínit například způsob eskalace po několikanásobném neúspěchu opravy. Při zapojení metod strojového učení se dále hovoří o tzv. self-healing systémech. Otázkou také zůstává, v jakých případech má být tato funkcionalita součástí sledovaného systému, či naopak systému dohledového.

3.5 Architektury dohledových systémů

Základním instrumentem pro analýzy příčin chyb ovšem zůstává kvalitní modelování architektury sledovaného systému souběžně s dohledovou částí. Z dobře popsané architektury primárního systému by mělo být možné poloautomaticky odvodit jak architekturu dohledového systému, tak detekční vzory. Ovšem co je skutečně možné a jak to efektivně dělat, je stále otevřeným problémem.

3.6 Analytické, návrhové a implementační vzory

Před cca deseti lety dospěl svět vývoje informačních systémů do stádia, kdy bylo již dost zrealizovaných analýz, návrhů a implementací, že bylo možné poznatky a vytvořené artefakty abstrahovat do podoby vzorů – analytických, návrhových a implementačních. Vzory jsou de facto opakovatelná schémata, postupy, znovu využitelné struktury napomáhající rychlejšímu postupu v jednotlivých etapách realizace systému při opakovaném řešení podobných problémů. Obdobná situace nastává nyní u řešení na bázi CEP, přičemž vzory jsou značně odlišné. Zavedení vzorů, jejich modelování (za použití i vhodné vizualizace), klasifikace a ověření si vyžadují seriózní výzkum.

3.7 Podpora aktivní role uživatele

Aplikace CEP, kam patří i dohledové systémy na bázi těchto technologií, lze zařadit na pomezí mezi tradičními řídicími a informačními systémy a systémy pro „business intelligence“. Do vývoje záhy vstupuje i koncový uživatel, aby rozšiřoval funkční systém o další pravidla (detekční metody), často jen za účelem jejich experimentálního ověření na aktuálních, historických nebo umělých testovacích datech. V ideálním případě by měl koncový uživatel či doménový expert mít možnost zadávat a zkoušet nová pravidla konverzací v přirozeném jazyce, což je nesnadné a v současnosti takřka nemožné implementovat. Přesto lze a je třeba uživateli dát prostřednictvím kvalitního uživatelského rozhraní sadu pomůcek, předpřipravených knihoven a „pískoviště“ (sandbox), pomocí nichž sestaví a ověří nové detekční metody.

3.8 Nedostatečná nativní podpora geografických informací

Dohledové systémy mají své nezastupitelné místo u správy budov, obecněji u tzv. *facility managementu*. Typické úlohy dohledu spočívají v korelaci událostí odehrávajících se velmi často nejen v blízkém čase (v určitém časovém okně), ale současně na geograficky blízkých místech – příkladem budiž současný vzrůst teploty v sousedních místnostech. Běžné systémy pro CEP nedisponují v sadě elementárních typů a operací prostředky pro snadné nalezení událostí, které souvisejí jak časově, tak prostorově. Je náročným výzkumně-vývojovým úkolem popsat a zkonstruovat taková rozšíření, jež budou umět tyto vztahy nejen zapsat a provádět, ale činit tak i dostatečně efektivně. Obdobné úlohy jsou časté i u jiných dohledových systémů – např. detekce podvodů při výběrech hotovosti ze vzájemně vzdálených bankomatů v rychlém sledu.

3.9 Interoperabilita dohledových systémů

S pokračujícím budováním dílčích dohledových systémů, ať už na bázi CEP nebo jiných, bude třeba stále častěji řešit souběžné fungování či integraci několika takových systémů budovaných potenciálně na různých platformách. Dodavatel dohledového řešení bude stále častěji v pozici integrátora než exkluzivního autora celého systému. Toto platí především pro distribuované systémy dohledu a jejich federace. Při výměně dat dohledu je třeba řešit řadu technických i konceptuálních potíží – různé protokoly výměny těchto dat, jejich formáty, sémantika, časové určení, z nichž jen některé lze vyřešit bez dalšího výzkumu a vývoje integračních technologií specifických pro dohledy.

3.10 Ochrana soukromí při dohledu

Při provozu dohledových systémů v počítačových infrastrukturách, sítích, ale i budovách se téměř vždy pohybujeme na hranici zasahování do práva na soukromí, případně i ochrany osobních údajů. Na druhou stranu, aby systémy dohledu plnily svůj účel, musejí určitá data o chování nejen strojů – systémů, sítí – ale i uživatelů zaznamenávat. To se musí díl v souladu s platnými předpisy, musí být především jasné, co, za jakými účelem, na jakou dobu je uchováváno a jak zpracováváno. Už jen správa pravidel a možnost jejich nezávislého ověření je významným výzkumně-vývojovým úkolem. Bude třeba výzkum, vývoj standardů a návazná příprava dobrých praktik tak, aby se poskytovatelé dohledových řešení měli při realizacích o co opřít.

3.11 Bezpečnost dohledových systémů

Souvisí s předchozím. Cílem je vyvinout koncepty, technologie a metody návrhu a nasazení zabezpečených dohledových systémů se zajištěním autorizace dohledu, zabezpečení přenosu atd. V kontextu dohledových systémů lze také očekávat manifestaci nových bezpečnostních hrozeb.

3.12 Forenzní využití dohledových systémů

Jednou z rolí dohledových systémů je odhalování protiprávního jednání. Aby data zachycená systémem byla akceptovatelná soudy jako důkazy, musí být prokazatelně zajištěna nejen jejich věcná správnost, ale i autenticita a integrita. Obdobně jako výše, musejí být vyvinuty standardy a sada dobrých praktik tak, aby se poskytovatelé dohledových řešení měli při realizacích o co opřít.

4. Shrnutí

Výše uvedené řešení i teprve identifikované problémy představují rozhodující výzkumně-vývojovou náplň laboratoře Lasaris na několik následujících let. Prostředí univerzity, nabízející jak V/V potenciál, tak příležitosti ověřovat navržená řešení v provozní praxi, skýtá ideální podmínky. Na svou příležitost čeká i případné komerční využití vytvořených konceptů a systémů formou spin-off společnosti či převzetím know-how jinou spolupracující firmou tak, jak se to již na Ústavu výpočetní techniky podařilo v případě detekce síťových hrozeb.

5. Literatura

- [1] Čeleda, P., Krejčí, R., Vykopal, J., Drašar, M. *Embedded Malware - An Analysis of the Chuck Norris Botnet*. In *European Conference on Computer Network Defense*. Vyd. 1. Los Alamitos, CA : IEEE Computer Society, 2010. ISBN 978-1-4244-9377-7, s. 3-10. 2010, Berlin, Germany.
- [2] Kučera, A. *Komplexní zpracování událostí v systémech pro správu budov*. Diplomová práce, Fakulta informatiky MU, 2012.
- [3] Luckham, D. *The Power of Events* Addison-Wesley Professional, 2002
- [4] Luckham, D. C., Frasca, B. *Complex Event Processing in Distributed Systems*
- [5] Nguyen, F., Pitner, T. *Scaling CEP to Infinity*. Sborník Letní školy aplikované informatiky, Bedřichov, 2012.
- [6] Nguyen, F., Pitner, T. Information System Monitoring and Notifications Using Complex Event Processing. In Zoran Budimac, Mirjana Ivanović, Miloš Radovanović (Eds.) *Fifth Balkan Conference in Informatics BCI 2012 Novi Sad, Serbia, September 16, 2012 Proceedings. Proceedings of the Fifth Balkan Conference in Informatics*. Serbia: ACM, 2012. od s. 211-216, 312 s. ISBN 978-1-4503-1240-0. doi:10.1145/2371316.2371358.
- [7] Pitner, T. Surveillance and Monitoring Systems Based on Complex Event Processing. In *15th International Conference on Business Information Systems*. 2012.
- [8] Tovarňák, D. *Monitoring rozsáhlé distribuované infrastruktury*. Sborník Letní školy aplikované informatiky, Bedřichov, 2012.

Ecosystem Condition Modeling Using Machine Learning Tools

Vadim Rukavitsyn

Department of Informatics, Faculty of Business and Economics, Mendel University
Zemědělská 1, 61300 Brno, Czech Republic
vadim.rukavitsyn@mendelu.cz

Abstract

There are a lot of different prediction methods exist in our days. Many of them work well in a specific situation. But when the situation deals with the big volume of data and complicated classification analysis there is a good opportunity to use the machine learning for the prediction. Previously there were a lot of different ways of applying the machine learning in the ecological modeling and prediction. But the method of the area quality estimation by the geofields analysis wasn't applied for the artificial intelligence. This method was proposed by Bondarenko and Zajonc and the classification was made only by the human expert. The method I proposed will develop this prediction method and it will add a new kind of the application of the machine learning in the environmental science. The method includes two parts. The first part is a special processing of the data of magnetic field, gravitation field and relief of the territory. This processing creates a base for the estimation of area quality. The second part is an applying of the machine learning tools for the final area estimation. This method will be very useful not only for ecological researches. Firstly, it will be very useful for the planning of new towns and for selection the place of residential areas building. Secondly, it will be useful for the territory estimation for real estate agencies. Thirdly, it will be helpful for the making of rational decision about the economical use of the territory.

Abstrakt

Existuje mnoho různých predikčních metod a mnohé z nich fungují i v konkrétní situaci. Ale když se situace řeší s velkým objemem dat, klasifikační analýza je dobrou příležitostí, jak využít strojového učení pro predikci. Dříve tam bylo mnoho různých způsobů, jak použití strojového učení v ekologickém modelování a predikci. Ale metoda odhadu kvality oblasti u analýzy geofields nebyla použita pro umělou inteligenci. Tato metoda byla navržena Bondarenkem a Zajoncem a klasifikace byla provedena pouze lidským odborníkem. Metoda, kterou jsem navrhl, bude rozvíjet tuto předpovědní metodu a přidá nový druh použití strojového učení v environmentální vědě. Metoda zahrnuje dvě části. První část je speciální zpracování dat magnetického pole, gravitačního pole a reliéfu území. Toto zpracování vytváří základnu pro odhad kvality plochy. Druhá část je použití tohoto nástroje strojového učení pro finální odhad oblasti. Tato metoda bude velmi užitečná nejen pro environmentální výzkum. Za prvé, bude to velmi užitečné pro plánování nových měst a pro výběr místa obytných oblastí budovy. Za druhé, bude užitečná pro realitní kanceláře pro území odhad. Za třetí, bude užitečná pro vytvoření racionální rozhodnutí o hospodárném využívání území.

Keywords

Machine learning, environmental modelling, ecosystem modelling

Klíčová slova

Strojové učení, environmentální modelování, modelování ekosystémů

1. Introduction

The main goal of this paper is:

- **Looking for the best method of the ecological division into districts automation according the ecosystem stability level.**

Most of all ecological division into districts was made by human experts. Especially if there is a complicated task and there is no way to use statistics for the classification making. In this situation there is only one way to automatize the process. This way is using of artificial intelligence tools. But in the same time the environmental data processing is a very hard task for this kind of analytical method.

ML involves a lot of different algorithms and programs. Task is that it is very important to choose the right one because different algorithms work good with one data and bad with another. Or in the same data different algorithms could show completely different significance. Also various programs can process the same data with different speed. The task is to find the most appropriate program and the best algorithm for data processing. Because in future analyzes it will save a lot of time and money.

The goal will be achieved by:

- **Applying machine-learning methods for the analysis of environmental datasets and predictions.**

Machine-learning could be applied in many situations and in different stages of work. First of all it is possible to estimate and sort data from datasets by ML, to solve a problem of missing values and so like. It can save a lot of time and money (if it will use some company). Then ML can choose some the most significant attributes of class estimation and make data processing more correct. After it data-mining methods could be applied to making classification of some new datasets. Finally ML applied for making the ecological prediction, and for the economic conclusion making for different situations.

2. Experiments and Their Results

Development of the geonatural systems in time and space could be recognized in the violation of the dynamic balance of the Earth and in the changing of lithosphere condition.

Because of those violations, the forming of the surplus stress structures happens. It causes the next relaxation of the stress, deformation and inner break of connections in the hierarchic chain of the geonatural environment organization. This process is a deformation of connections in ultimate particles. It is accompanied by the forming of new molecular structures and chemical compounds. Then it changes the structure of geological and natural systems.

Each changing of the energetic field which cause the violation in ultimate particles connections, produce anomalies in the magnetic field and gravitation field. Therefore each geofield can show us chemical and structural features of the geological environment.

It is obvious that every natural object could be described by the combination of the physical parameters. For example Earth surface could be characterized by space photos, maps of magnetic and gravitation fields, relief maps and so on. Each from those features describes geosphere in its own hierarchic level. Mathematical transformations of the physical field combination allow us to get the picture of processes dynamic on the area. It could help us to understand process of the environment development and connections between geofields changes and natural processes.

Heterogeneity of every physical field of the Earth is a derivative of integral changes of its stress condition. The calculation of primary components allows us to get a vector field and to create the map of the static condition of geological and natural systems.

Applying of spline methods allow us to find the most important anomalies of the system condition.

Geological potential of some geological environment shows us its potential space energy. Every energy process (endogenous and exogenous processes too) goes with potential energy loss in time of the transition from the area on maximum stress to the area of minimum stress. Therefore dynamic models allow us to see not only the potential energy of the geofields deformation but also to recognize their orientation.

For the mathematic description of the energy processes speed divergence calculation is used. It could be the function from the coordinates and time. It is the same for endogenous and exogenous processes as the speed for the mechanic move. Divergence shows us the activity and the duration of processes.

By the combination of geofields analysis methods, mathematical transformations and processes dynamic estimation it is possible to create the forecast model of the geonatural system condition. This model allows us:

1. to show different conditions of the geonatural system forming,
2. to investigate the impact of system changes on the biosphere and other environment components,
3. to combine all the variety of geological and natural factors in systems,
4. to model the processes of the geonatural system changes and describe the dynamic of environmental components iteration modes.
5. to allocate structural and functional heterogeneity of the environment.
6. to get the picture of the environmental processes dynamics.

The scale of the original data determines the accuracy of the prediction model. But the combine processing of the differently scaled data of geofields and relief geodynamic transformations could make the forecast more precise.

Ecological division of an area into districts is a method that estimates dangerous genetic changes of the environment. It could also recognize problem situations, areas, spheres and economic objects which make the biggest influence on negative changes of an ecosystem and population health.

Results of this kind of the ecological classification, expert estimations and support of decision making that connected with the decreasing of the industrial impact on the environment are the complex of ecological measures which are necessary for the perspective planning of city structures and region structures. Those measures help with the distribution of industrial, transport and social objects on the given territory.

This work is an attempt to create of the analytic basic model by the machine learning tools. This model has to classify territories, estimates the level of ecologically dangerous changes of the environment and recognizes problem situations and areas of their localization. The model has to show the heterogeneity of geodynamic conditions of the territory forming and it has to estimate the level of the ecosystem condition.

Existed nonlinear thermodynamic connections between the modern temperature regime of the lithosphere and the activity of chemical and biological processes determine different speeds of geochemical and biochemical reactions, time of live organisms development, different intensity of photosynthesis etc.

There are three types of geodynamic processes which determine an ecosystem condition changes and the level of the ecosystem stability (ecosystem resistibility and ability to restore violated connections and functions). Those types are: stable, normal and unstable.

Stable geodynamic fields are characterized by the long geological environment consolidation in consequence of the thermal streams activity decreasing and the earth surface subsidence. The reduced thermal regimes level and the stable geodynamic situation determine conditions of the longer organism's development, the reducing of the geochemical and biochemical reactions speed, the rise of the photosynthesis production. All those factors stimulate the rise of the environment resistance to man-caused stress level. In stable geodynamic zones the environment can restore its violated connections and functions. This feature is necessary to take into account in the spa-zones creation, in the medical complexes placement and in the production of bio-products. Also the stable ecosystem condition cause low morbidity level in this zone. It means that the morbidity level in the region could be a great tester of the model accuracy.

Unstable geodynamic processes are characterized by the fast softening of the environment in consequence of thermal, radioactive and other impacts. The increased thermal regime of the geonatural system, the geochemical and biochemical reaction activity rise cause the development of

destabilization and environment destruction processes. In those zones there is the biggest possibility for appearance of ecologically dangerous processes. There are processes like a rise of soil and water pollution activity, an activation of modern tectonic processes. This activation of modern tectonic processes could cause an appearance and development of landslides, underflooding processes, destruction of engineer constructions and communications.

There are zones of the combined natural and anthropogenic impact on the population health deterioration.

It is necessary to mention that ecologically dangerous genetic changes of the ecosystem and the population health deterioration are caused by the complex of natural and anthropogenic factors. Geodynamic condition of the ecosystem forming is a very important but it is not the only one factor which influence on the environmental situation. There are four main groups of factors:

1. Natural processes and phenomenon. There are a non-artificial event in the physical sense, and therefore not produced by humans. Common examples of natural phenomena include volcanic eruptions, weather, decay, gravity and erosion. This group of factors has direct and obvious influence on population and ecosystem.
2. Anthropogenic factors. There is the direct human influence on the ecosystem. The anthropogenic impact includes impacts on biophysical environments, biodiversity and other resources (Sahney S. 2010). Human influences almost on every part of the ecosystem. Most of all this influence is conditioned on industries (Agriculture, energy industry, manufactured products, mining, transport, war etc.).
3. Socio-demographic factors. This group of factors influences mostly on the population and determines a level of life, an amount of resources per person etc. It clearly connected with the morbidity level of the population. A good example of this connection is social level and people health. A good start in life means supporting mothers and young children: the health impact of early development and education lasts a lifetime. Observational research and intervention studies show that the foundations of adult health are laid in early childhood and before birth. Slow growth and poor emotional support raise the lifetime risk of poor physical health and reduce physical, cognitive and emotional functioning in adulthood. Poor early experience and slow growth become embedded in biology during the processes of development, and form the basis of the individual's biological and human capital, which affects health throughout life. Infant experience is important to later health because of the continued malleability of biological systems. As cognitive, emotional and sensory inputs program the brain's responses, insecure emotional attachment and poor stimulation can lead to reduced readiness for school, low educational attainment, and problem behavior, and the risk of social marginalization in adulthood. Good health-related habits, such as eating sensibly, exercising and not smoking, are associated with parental and peer group examples, and with good education. Slow or retarded physical growth in infancy is associated with reduced cardiovascular, respiratory, pancreatic and kidney development and function, which increase the risk of illness in adulthood.
Poverty, relative deprivation and social exclusion have a major impact on health and premature death, and the chances of living in poverty are loaded heavily against some social groups. Absolute poverty – a lack of the basic material necessities of life – continues to exist, even in the richest countries of Europe. The unemployed, many ethnic minority groups, guest workers, disabled people, refugees and homeless people are at particular risk. Those living on the streets suffer the highest rates of premature death (Marmot Michael. 2003)
4. Geodynamic conditions of the ecosystem forming. This factor influences on the whole ecosystem and population in the same time and its influence is high. It has the same and sometimes even higher importance than all mentioned factors. But the impact degree of the geodynamic conditions is not the same in different zones. This degree depends on the types and activity of geodynamic fields. If there are active stable or active unstable fields, the degree of the influence and the factor importance will increase. In normal or not active zones the geodynamic fields give about 40% of the whole impact on the ecosystem but in active zones

the there could be 60% of the whole impact. In any case this factor is very important and couldn't be ignored (Zayonts I. O. 2001).

3. Modeling

There is an attempt to create the analytic model that bases on the geodynamic fields analysis. In the machine-learning processing it will be shown only one application of this prediction model. I will create and test the model by the possibility to find the impact between the geonatural system changes and the biosphere condition. Previously all described applications of the method were made only by the human experts and the using of machine-learning tools is the best way to automatize the process.

For the first experiment and for the creation of the model there was used the data from Kiev. It was necessary to create the model which was based on the train data. There was the classified data of the magnetic field, gravitation field and relief of the city and it contained 193 830 points. The first classification was made by human experts. Experts divided the territory to different zones by the environmental system conditions. Zones of the stable environmental system condition were corresponded to stable biosphere condition zones. Normal environmental system zones were corresponded to normal biosphere condition zones and unstable environment system complied with the unstable biosphere condition.

Every machine-learning processing needs previous data preparations. It means that is necessary to chose right attributes. An attribute is a property of an instance that may be used to determine its classification. Determining useful attributes that can be reasonably calculated may be a difficult job (<http://www.cse.unsw.edu.au/~billw/mldict.html>).

As the first three attributes I have chosen the magnetic field, the gravitation field and the relief. Those three parameters were mathematically transformed for the creation of the rest of attributes.

There is the list of used attributes:

1. Magnetic field
2. Gravitation field
3. Relief
4. The first principal component
5. The second principal component
6. Divergence based on the first principal component data
7. Divergence based on the second principal component data
8. Spline of the first principal component data
9. Spline of the second principal component data
10. Divergence based on the data of the first principal component data spline
11. Divergence based on the data of the second principal component data spline
12. Class

In the classification I used 5 classes of the territory: a very unstable zone, an unstable zone, a normal zone, a stable zone, a very stable zone.

Experiments described below were made by WEKA software and See5 software. For the searching of the algorithm there was used 10-folds cross validation. In those experiments 16 algorithms were used. For each of those algorithms the best parameters were chosen. In the table below you can see results of the best parameters searching.

Table. 1. The best algorithms with the best parameters.

Algorithms	Result (%)	Best parameters
J48	97,12	confidenceFactor 0,01; minNumObj 2; numFolds2; seed 1
Random Forest	98,25	maxDepth 0; numFeatures 0; numTrees 12; seed 1
Random Tree	97,28	KValue 4; maxDepth 0; minNum 1; num Folds 0; seed 5
Simple Cart	97,00	minNumObj 2; numFoldsPruning 5; seed 3; sizePer 0,9

See5	98,10	Amount of boosts 10
IBk	96,90	KNN 5; nearestNeighbourSearchAlgorithm BallTree; windowSize 0
Bagging	97,57	bagSizePercent 100; classifier Random Forest; numIterations 10; seed 1
Decorate	97,49	artificialSize 1,0; classifier J48; desiredSize 20; numIterations 10; seed 1
END	97,56	classifier J48; numIterations 10; seed 1
Classification Via Regression	97,25	Classifier Random Tree
Data Near Balanced ND	97,16	classifier Random Fores; seed 2
ND	97,16	classifier J48; seed 1
Class Balanced ND	97,14	classifier Random Fores; seed 5
Random Subspace	97,56	classifier J48; numIterations 10; seed 1; subSpaceSize 0,5.
Rotation Forest	97,20	classifier J48; maxGroup 3; minGroup 3; numIterations 10; projectionFilter PrincipalComponents; removedPercentage 50; seed 1

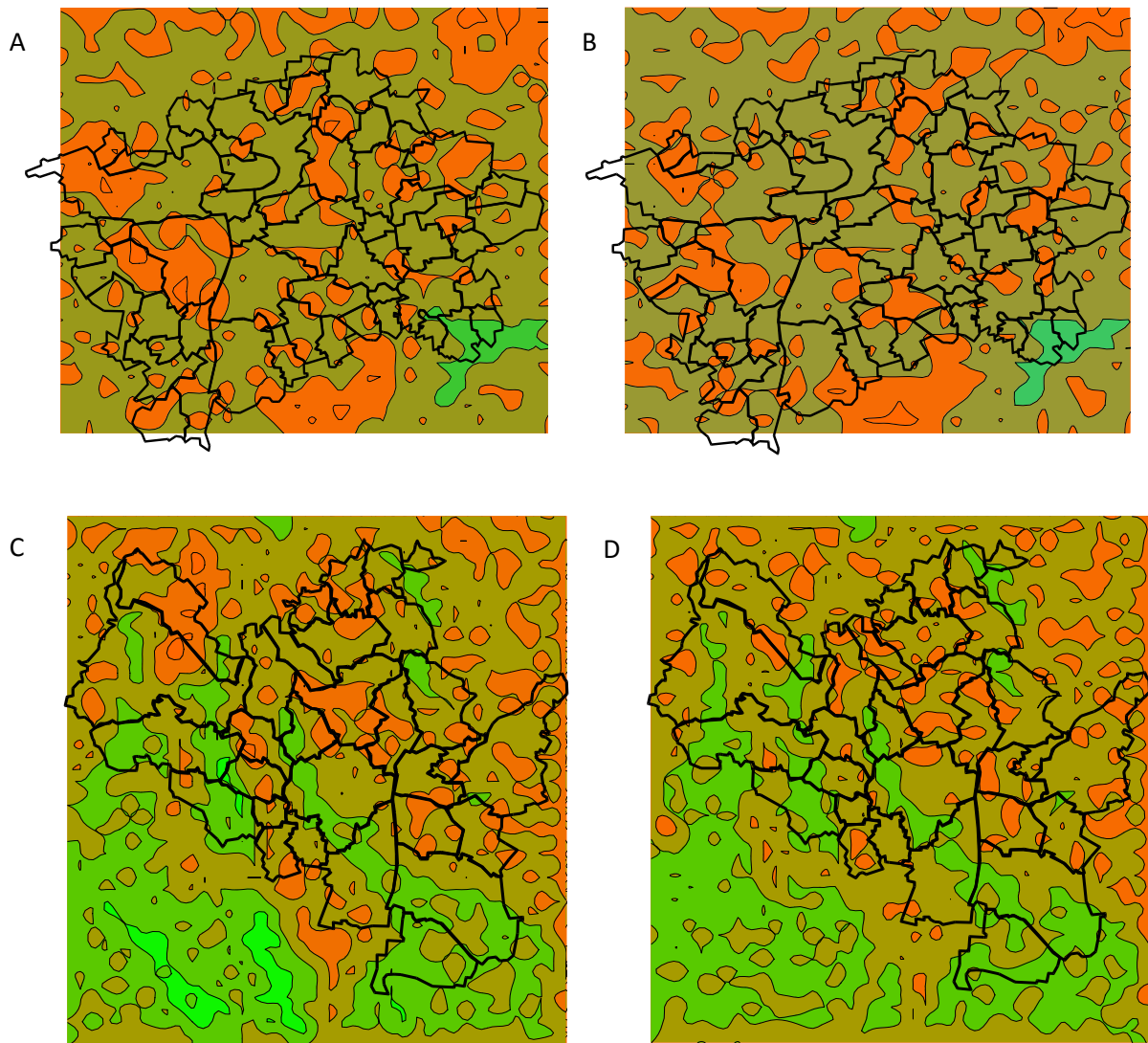
As you can see the best result was provided by the algorithm Random Forest with more than 98% of accuracy. It means that the computer model could copy the logic of the human expert in this kind of the classification. But here is one of the main problems of the environmental data processing by the machine learning tools. Even if the computer analyzed one situation and cross-validation showed good results, in another similar situation the result could be worse. It happens because every nature situation is unique and it is very difficult to create an absolutely adequate model. So big model accuracy in the training dataset shows us that the computer copied an expert logic in the classification of this concrete territory with this concrete conditions.

4. Prediction

For the task of prediction I used 3 different datasets. There were data from all Russia territory, from Prague and Brno. Now the data was not classified by a human expert. Therefore it was an excellent way to test the model in the real work. All the estimation on the testing set should be made only by a computer and then results could be analyzed separately. The thing is that the method of the ecosystem estimation should be universal and work equally well with different scales. Exactly for the testing of its universality so different datasets were chosen. The fact that those datasets are from different parts of the world will approve the significance of the method even better. As an estimator of the prediction accuracy I used activity of the morbidity level in different regions of the territory. It is a good indicator of the ecosystem condition, especially if it deals with congenital diseases because they have the biggest correlation with the environment and don't depend so much from the lifestyle. In the end there is possible to compare the predicted activities with the real situation and estimate the accuracy of the model.

Firstly the method was applies on Prague and Brno regions. Before the classification the data from the area was collected and prepared. Original data contains information about the magnetic field, gravitation field and relief on those territories (CGS. 2012). The data was collected form the Czech geological service. Brno data contained 7209 points and Prague data contained 17061 points. Before the attribute calculation there was necessary to create the mutual grin of coordinates because all parts of the original data have different probe points on the same territory.

After the preparation attributes were calculated and the modeling started. The training dataset from Kiev was used like the base of the model and datasets from Prague and Brno were the testing datasets. For the modeling two algorithms were used. There were Random Forest and Bagging algorithms. Fandom Forest showed the best accuracy and Bagging had to control the Random Forest because it had the best results not from the decision trees algorithms. The results you can see below.



Class of the ecosystem condition:

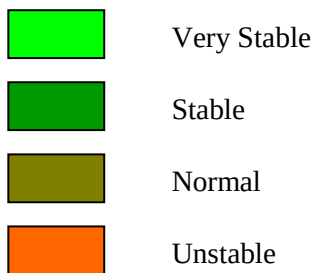


Figure1. Model of the ecosystem condition. A) Prague. Random forest algorithm; B) Prague. Bagging algorithm; C) Brno. Random forest algorithm; D) Brno. Bagging algorithm.

As you can see in there are more stable areas the Brno territory. Some districts in Brno contain even very stable territories. The conformation of this modeling is in the medical statistics which show that in Prague there are around 450 congenital diseases per every 10 000 people and in Brno there are less than 260 (UZIS. 2011). There were count all congenital diseases from 2005 till 2009. According those materials even if the absolute amount of diseases changes the relative amount stills almost the same. It helps us to make a conclusion that there are some long term circumstances, which cause this situation. Also there was build the map of the congenital diseases in Prague. You can find it on the figure 2.

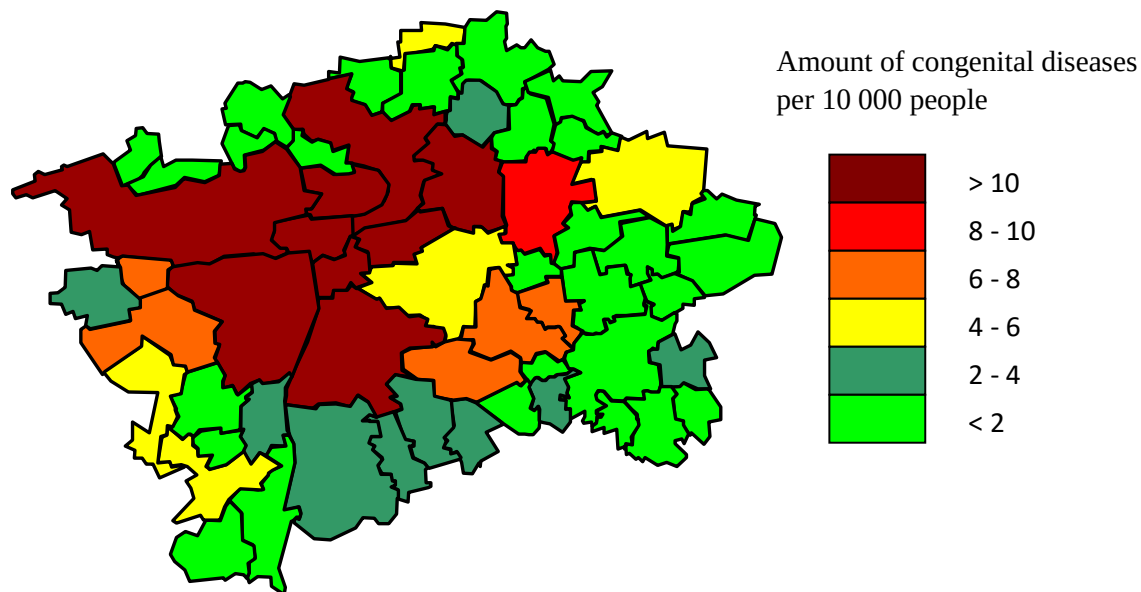


Figure 2. Congenital diseases in Prague per every 10 000 people.

According to the model, in all the Prague territory is almost the same situation. It is possible to mark out only a one big stable region on the south-east and several relatively big unstable regions on the west, north and in the center. On the map of diseases you could see that the biggest amount of diseases is exactly in the unstable regions and in all south-east part is a relatively good situation and the smallest amount of diseases. It shows that the model corresponds with the real situation and can make a significant prediction.

For one more model testing there was used data of another scale and another territory. There was dataset which was based on the information from Russia.

For the building of Russia territory model the data was collected from maps. There were maps gravitation (Demyanov G.V. et. al. 1995) and magnetic fields (Litvinova T.P. et. al. 1995) and a map of Russia relief (Lagutina N.P. et. al. 2005). Those maps were digitized and all the information was collected in one table. After it, this data, which contained the information about two geofields and relief, was processed. There were calculated 2 primary components, splines of those components and the divergence.

Then a classification was made which was based in the dataset from Kiev. It means that Kiev dataset was the training data and the data from Russia was the testing data. It contained 5900 points. For the building of the model I used 2 algorithms. The first one is Bagging because it showed the best results from all algorithms not from the decision tree group. The second one is Random Forest because it had the best accuracy. On the pictures below you can see results of the classifications.

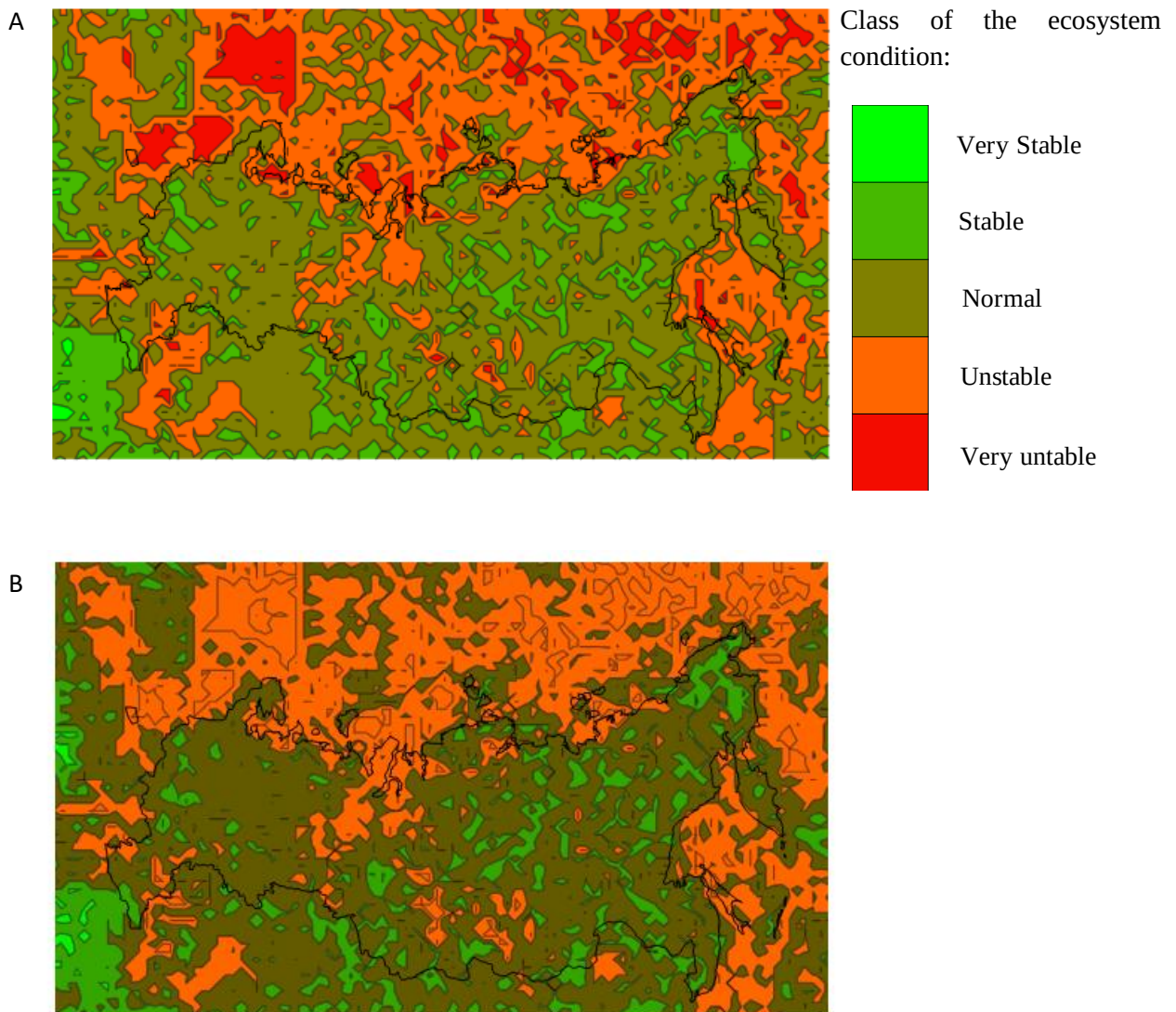
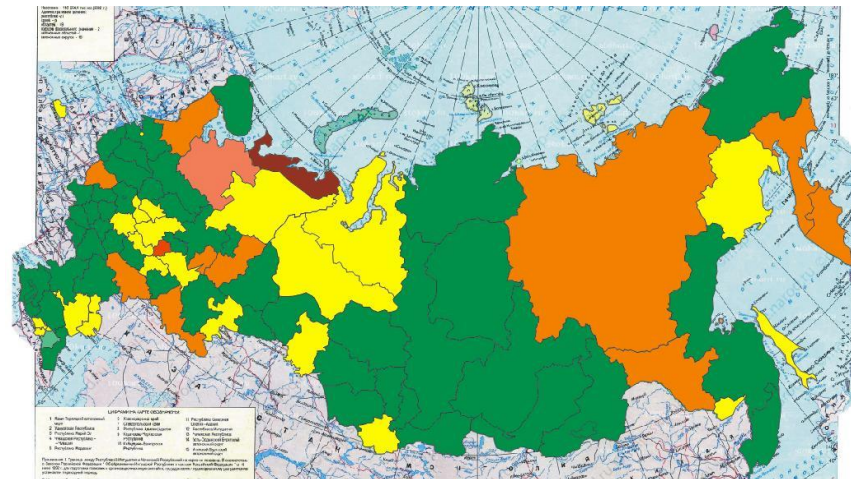


Figure 3. Ecosystem condition in Russia estimated by: A) Bagging algorithm; B) Random forest algorithm

Both classifications are very similar but the random forest algorithm didn't recognize the very unstable class of the territory. Now it is necessary to compare those models with the morbidity in Russia for making a final conclusion. Unfortunately it is possible to estimate the morbidity only according the regions. It is almost impossible to create an authentic isoline map of morbidity because all the information is sorted by region. The most reliable information will be about a western part of Russia because regions there are relatively small and the situation is better shown.

For making the morbidity map there have been used Russian medical statistical materials (Ministry of Health and Social Development of the Russian Federation 2011).

I collected the information about congenital diseases in every region. There was used the relative morbidity. It means that it was estimated an amount of congenital diseases by 100 000 people. There were calculated only congenital diseases and other was not. It was like that, because congenital diseases could be caused by the environmental situation with bigger possibility than other diseases, which could be caused by the lifestyle or the climate. There were calculated congenital disorders, congenital blood defects, congenital nervous system defects and congenital gynecological defects. On the figure 4 you can see the morbidity map.



Amount of congenital diseases per 100 000 people:

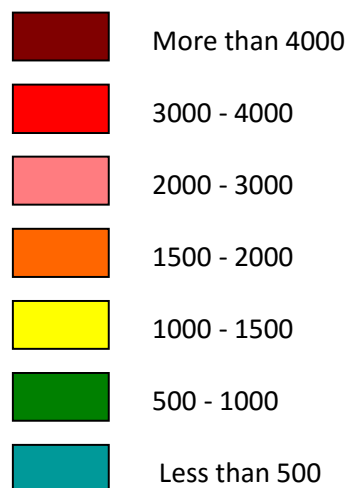


Figure 4. Morbidity in Russia

It is possible to say that the correlation between the congenital morbidity and the machine-learning territory estimation is very big. Especially it is good shown in the west part of Russia where regions were not too big and well populated in the same time. Here the northern part has the biggest amount of diseases and is situated in the unstable and very unstable zones. In the same time Kavkaz is the healthiest part of the country and situated in the stable and very stable zones. Bagging algorithm has shown itself better than Random Forest. The model describes here the general view on the situation and shows local centers of negative and positive processes. According this information I can consider that the computer model works well on the test data. Also I can consider that the model is universal because I trained this model on the dataset from one city and tested it on the dataset from whole country. Despite the hard task the testing was quite successful.

5. Conclusion

There are a lot of different prediction methods exist in our days. Many of them work well in a specific situation. But when the situation deals with the big volume of data and complicated classification analysis there is a good opportunity to use the machine learning for the prediction. Previously there were a lot of different ways of the machine learning applying in the ecological modeling and prediction. But the method of the ecosystem condition modeling by the geofields analysis wasn't applied for the artificial intelligence. Previously such kind of the classification was made only by the

human expert. The described method automatizes expert analysis of this problem and adds a new kind of the application of the machine learning in the environmental science.

The goal of the work was looking for the best method of the automation of ecological division into districts by the ecosystem stability level. The method was found. It combines ways of the data preprocessing for the solving of this problem and the using of concrete algorithms for the best modeling of the process. But the significance of the model in testing was lower than in training. It happened because of shortages of the training dataset. The significance around 80% in different testing datasets showed that the model works and all the preparations and the algorithm choose were right. But the analysis of the model significance in the estimation of separate classes let us see that the problem was in unbalanced training data. Also the model was based on the data from one small region. When it deals with the ecological modeling it is better to collect the information about different standard objects from all over the world. If we create a big balanced dataset which contains the information about many different areas, the model accuracy will be close to 98%.

The fact that the model could make the estimation of the areas with different scales and from different regions tells us that the method is universal. The universality determines by three parameters. The technology should work in any scales, in any territories and any time. The method of ecological division into districts by the ecosystem stability level meets all three terms and it was shown in experiments.

Automation was working well too. By the using of the machine learning the estimation could be much faster and easier. If we combine the described method with the information system, we will create the decision support system which can help in the solving of many environmental problems. Of course the role of the expert will be still very important and couldn't be totally substituted by the computer but his work will be faster and more effective with using of the showed technology in the ecosystem analysis.

This method will be very useful not only for ecological researches. Firstly, it will be very useful for the planning of new towns and for selection the place of residential areas building. Secondly, it will be useful for the territory estimation for real estate agencies. Thirdly, it will be helpful for the making of rational decisions about the economical use of the territory.

Also it is important to mention that the described method build the ecosystem model by the very significant but not the only one factor what can influence the environment. If we add to the big balanced training dataset the information about natural phenomena, anthropogenic factors and socio-demographic conditions, the machine-learning model will be not only precise but also it will create the complex view on the ecological situation in the studied area.

6. References

- [1] Demyanov G.V., Nazarova N.G., Nikolsky Yu.I., Taranova V.A. (1995). *Map of Anomalous Gravitational Fields of Russia and Adjoining Water Covered Areas*. Scale 1:10 000 000. VIRG-Rudgeofizika.
- [2] Czech Geological Service (2012). Magnitometric and gravimetric data from Prague and Brno region.
- [3] Lagutina N.P., Ostovskaya O.E. (2005). *Map of Relief of Russia*. Scale 1:10 000 000. FGUP «Kartographyc factory of Omsk»
- [4] Litvinova T.P., Shmiyarova N.P. (1995). *Map of Anomalous Magnetic Field of Russia and Adjacent Water Areas*. Scale 1:10 000 000. VSEGEI.
- [5] Marmot Michael, Richard Wilkinson (2003). *Social Determinants of Health: The Solid Fact (Second Edition)*. World Health Organization (Europe).
- [6] Ministry of Health and Social Development of the Russian Federation (2011). *Morbidity of the Population in Russia in 2010. Statistical materials. Part 2*. Moscow. pp. 114-118
- [7] Sahney, S., Benton, M.J., Ferry, P.A. (2010). *Links between global taxonomic diversity, ecological diversity and the expansion of vertebrates on land*. Biology Letters 6. pp 544-547.
- [8] UZIS CR (2011). *Congenital anomalies in births in year 2009*. Translation UZIS CR.
- [9] Zayonts I. O., Bondarenko J. J., Slipchenko B., Lysychenko G. V. (2001). *New approaches to the problem of geoecological risk for urbanised territories*. ECO-INFORMA 2001. Chicago, USA.

Indoor Navigation for Mobile Devices

Jonáš Ševčík

Masarykova univerzita, Fakulta informatiky
Botanická 68a, 602 00 Brno, Czech Republic
255493@mail.muni.cz

Abstract

This article deals with techniques suitable for indoor navigation using mobile devices. It presents several principles applicable to mobile devices. Followingly, these principles are used in demonstrative Android application, which combines them into working indoor navigation prototype. Consequently, demo application is used for gathering accurate results of the above mentioned principles. Obtained results are presented in the last section of this paper.

Abstrakt

Příspěvek se zabývá technikami vhodnými pro navigaci pomocí mobilních zařízení v uzavřených prostorech a představuje několik řešení vhodných pro tato zařízení. Tato řešení jsou následně využita v ukázkové aplikaci pro platformu Android. Vytvořená aplikace kombinuje jednotlivé způsoby navigace ve funkční navigační prototyp. Tato aplikace je poté využita ke zjištění přesnosti výše zmíněných způsobů navigace. Získané výsledky jsou zaznamenány na konci tohoto článku.

Keywords

Indoor navigation, Android, Wi-Fi localization, dead reckoning, particle filtering, step detection.

Klíčová slova

Navigace v uzavřených prostorech, Android, Wi-Fi lokalizace, dead reckoning, částicové filtrování, detekce kroků.

1. Introduction

In our previously hold research there was a need to build an indoor navigation solution for mobile devices [9]. Therefore our team has started implementing prototype navigation for Android platform. Nevertheless solutions used while implementing this prototype are generally applicable to any other operating system.

According to the articles [8, 9] we are not able to use Global Positioning System (GPS) in closed environments. Therefore we are limited to the use of sensors present in mobile device. From all the possible sensors we have selected following ones; Compass/magnetic field sensor, gyroscope and accelerometer. Also, we took in consideration radio wave based location techniques, therefore we include to this list Wi-Fi adapter as well. These listed pieces of hardware can be part of the equipment of Android powered cellular phone.

2. Localization Methods

In this section we would like to introduce a brief overview of available localization techniques suitable for mobile use.

2.1 Dead Reckoning

Dead reckoning is a relative localization method, which uses calculation as a mean of obtaining position. Calculation is based on data acquired by measuring motion and direction change from initial position.

As a special type of dead reckoning navigation technique, we can distinguish so called inertial navigation. In this technique, the position is computed by double integration of input obtained from accelerometer.

In this article we are going to present dead reckoning in the sense of inertial navigation.

2.2 Electromagnetic Wave Based Localization

These techniques contain previously mentioned GPS or other similar systems like Galileo and Glonass. Thus more suitable candidates are radio wave based methods using for example Wi-Fi, Bluetooth, Radio-frequency identification (RFID) etc.

Location estimation is calculated by comparing Received Signal Strength Indication (RSSI) to a broadcast reference signal strength level. In the case of RFID, there is a requirement for a prepared environment containing evenly spread RFID tags.

2.3 Image Processing

Image processing methods can be based on feature or marker detection. In the marker based detection, an image of surrounding environment is compared with marker database. This database contains unique markers mapped to their coordinates in the real world. In the case that known marker is discovered by scanning device, its coordinates are obtained from previously prepared database.

3. Used Techniques

In this section we are going to introduce techniques we implemented in the prototype application.

3.1 Dead Reckoning

In order to use this technique it was necessary to resolve how to measure motion. We decided to use step detection. As presented in [3, 4, 6] currently used procedures were based on data gathered from an accelerometer unit placed on a suitable part of a human body e.g. foot or hip. Our approach cannot use such device placement because we need the user to watch the screen of the device while walking. The device must be held in hands.

The following list contains digital signal processing filters [7] suitable for step detection;

- Low-pass filter – removes high frequency components of the signal. This is important to avoid false positive steps.
- Power threshold filter – computes signal power, which is used for identification of void zones.
- Step duration change filter – It is probable that duration of two succeeding steps does not change distinctly. Detected step candidate can be rejected when its duration is less than half the duration of the previous step.
- Correlation comparison filter – computes similarity of two signals by their correlation.

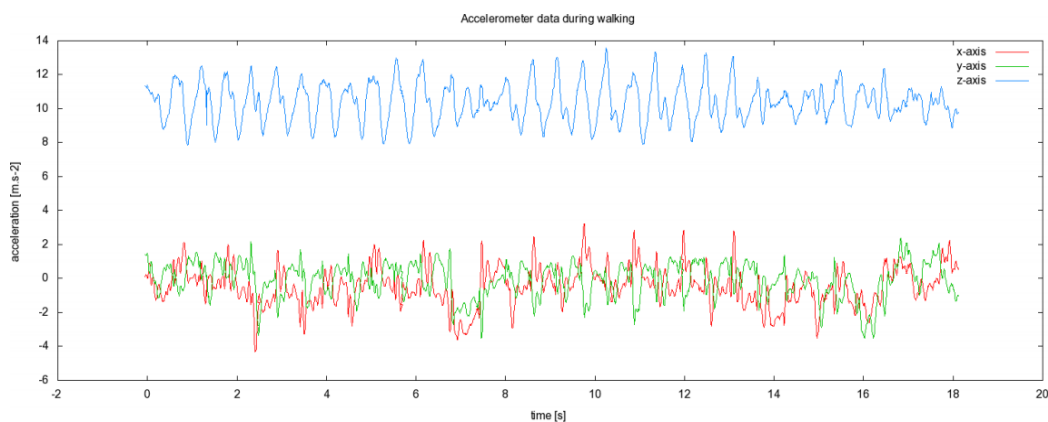


Figure 25: Accelerometer signal on the z-axis; Steps measured at 100 Hz sampling rate.

We implemented step detection as a combination of two moving average filters with different window sizes. The longer moving average (window size of 0.2 s, which is an equivalent of twenty samples at 100 Hz sampling rate) is used as the average value of g-force estimation projection on the z-axis. The shorter moving average (windows sizes of 0.05 s; five samples; 100 Hz sampling rate) is a low-pass filter that removes the high frequency components from the signal. Candidate steps are detected as intersections of two moving average signals.

The moving average method proved to be successful in tests and was correctly detecting steps with near 100% precision during sustained walking. Problems were detected when the detector detected false positive steps while standing still.

To eliminate this flaw, the algorithm was extended by adding a power threshold filter, that defines the walking and non-walking zones in the signal.

The algorithm may still detect false steps due to acceleration signal noise caused by manipulation with the held device. We decided to handle false positive steps by successive mathematical processing. Power threshold has to be configured manually; no self-adjusting algorithm was used [5].

3.2 Wi-Fi Localization

We decided to use Received Signal Strength (RSS) fingerprinting method over standard triangulation, because we did not have the information about the location of access points.

RSS fingerprinting consist of two phases. Offline phase consists of fingerprint gathering. During this phase the database of RSS fingerprints is created. Fingerprint is symbolized as a pair (x, y) , where x are a coordinates in the real environment; y is a unique access point identifier and the strength of received signal measured in dBm.

During the online phase, the measured signal strength along with access point identifier are compared against the database of gathered fingerprints. The best matching candidate is chosen as an estimated location provider.

The phase of gathering fingerprints can be simplified as shown in [2]. This method consists of gathering a smaller amount of fingerprints and interpolating the gathered data to obtain the full map image. Measured map is triangulated using Delaunay algorithm.

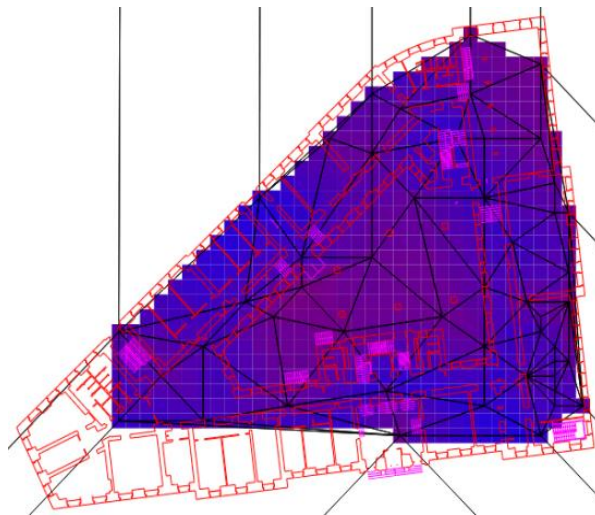


Figure 26: Interpolated RSS map.

The triangulation algorithm splits the area into non-overlapping triangle mesh, where each location in the plane is assigned three unique vertices of the triangle. Full linear interpolation of the map can be recreated from the Barycentric coordinates of the location in the triangle.

3.3 Sequential Monte Carlo Filtering

Sequential Monte Carlo (SMC) filtering is modeling the state of a dynamic system by approximating the posterior density function by a set of random samples of the state vector while sequences of noisy measurements are made on the system.

The computation of the SMC requires two models [1]:

1. System model – model describing the evolution of the state and time;
2. Measurement model – model relating the noisy measurements to the state.

Filter processes events from the step detector (step event, length estimate) and a probability density function from the Wi-Fi AP scans.

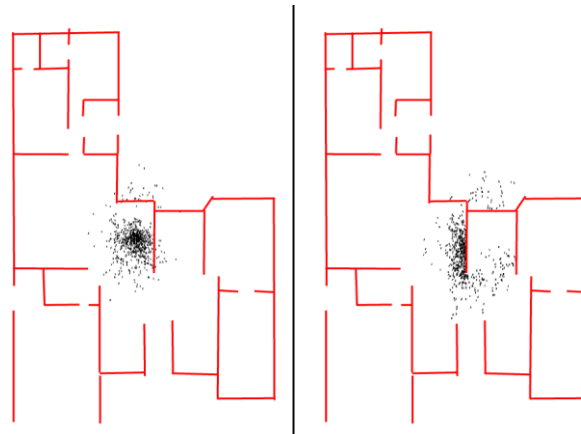


Figure 27: Elimination of particles on impassable obstacle.

As shown in [1], SMC uses particles which are placed and evenly spread in the probable location determined by RSS fingerprint database. These particles are set to motion with events generated by step detection. Followingly, those particles which hypothetical motion leads through impassable obstacles, e.g. walls, are eliminated. As shown in figure 3. This results in improvement of location estimation.

4. Application

We have developed demo application for Android platform to measure accuracy of before mentioned techniques. Main part of the application is called Pedestrian Localization Prototype.



Figure 28: Indoor Localization Prototype.

In the Figure 4 there is clearly visible blue rectangle area which is the representation of estimated location obtained via Wi-Fi localization. Inside, there is a blue circle representing a user and evenly spread particles. Once the user starts moving, particles help to calibrate the estimation of his location.

5. Precision Testing

We had randomly selected a set of points on the ground floor of Faculty of Social Studies building. Each point was assigned real coordinates. Followingly, using Samsung Galaxy Tab as a testing device we performed a random walk on the previously selected points. After reaching a point, we acquired its estimated location and compared it to the actual location. We made two precision tests. First test was performed with known starting position, using dead reckoning and particle filter. Second test was started without known starting position, using dead reckoning, particle filter and Wi-Fi localization. We made a total of 75 measurements. During the first test, two measurements could not be successfully completed due to a position loss and one measurement had a distance error of 66 meters, thus these measurements were discarded from the overall score.

First test finished with measured median distance error of 2.3 meters, 90th-percentile of 5.6 meters. Second test results are median distance error of 3.6 meters and 90th-percentile of 12.7 meters.

6. Conclusion

We have researched suitable indoor location techniques from which we have selected dead reckoning, particle filtering and Wi-Fi localization as the most suitable ones. These techniques were used as main techniques in prototype Android application. This prototype was used to test precision of dead reckoning, particle filter with and without Wi-Fi localization. Test results show that Wi-Fi localization technique did not have desired improving effect on dead reckoning position estimate. Particle filter had difficulties converging to the correct location after estimating initial position using Wi-Fi localization. According to [8] it is recommended to use AP topology rich in overlapping. Low precision of Wi-Fi localization may be caused by optimal distribution of AP topology.

6.1 Future Work

We would like to improve precision of Wi-Fi localization along with the particle filter. Achieved precision error of 2.3 meters is allowable for successful localization within building corridors. Therefore we would like to use implemented techniques as a base of an indoor navigation application for the buildings of Masaryk University.

7. Acknowledgements

We would like to thank all the staff members of the Masaryk University Department of Building Passportization. They provided maps and consultations needed for successful implementation of our prototype.

8. References

- [1] Arulampalam, M. and Maskell, S. and Gordon, N. and Clapp, T.: A Tutorial on Particle Filters for Online Nonlinear/Non-Gaussian Bayesian Tracking, IEEE Transactions on Signal Processing, Vol. 50, No. 2, February 2002.
- [2] Chen, A. and Harko, C. and Lambert, D. and Whiting, P.: An Algorithm for Fast, Model-Free Tracking Indoors, Mobile Computing and Communications Review, Volume 11, Number 3, 2006.
- [3] Cho, S.: MEMS Based Pedestrian Navigation System, The Journal of Navigation (2006), 59, pp. 135–153, The Royal Institute of Navigation, 2006.
- [4] Frank, K. and Krach, B. and Catterall, N. and Robertson, P.: Development and Evaluation of a Combined WLAN & Inertial Indoor Pedestrian Positioning System, 4th International Symposium on Location and Context Awareness, ION GNSS, Savannah, Georgia, USA, 2009.
- [5] Holčík, M.: Indoor Navigation for Android, Master's thesis, Masaryk University, Brno, 2012.
- [6] Libby, R.: A simple method for reliable footstep detection on embedded sensor platforms, 2008.
- [7] Smith, S.: The Scientist and Engineer's Guide to Digital Signal Processing, 1997-2006.
- [8] Ševčík, J. and Tokárová, L: Application of Augmented Reality on Mobile Devices, 8. Letní škola aplikované informatiky, MSD Brno, 2011.
- [9] Ševčík, J.: Rozšířená realita na mobilní platformě Android a její aplikace v knihovnictví, Master's thesis, Masaryk University, Brno, 2011

Princípy doručovania BI na mobilné zariadenia

Lucia Tokárová

Masarykova univerzita, Fakulta informatiky
Botanická 68a, 602 00 Brno, Czech Republic
173309@mail.muni.cz

Abstrakt

Tento článok sa zaoberá konceptom mobilnej business inteligencie (BI) z pohľadu užívateľskej prívetivosti. Cieľom štúdie bolo vytvorenie sady princípov doručovania BI na mobilné zariadenia. Prvým krokom bola formulácia poznatkov o užívateľoch a kontexte použitia mobilných BI nástrojov. Ďalej nasledovala štrukturovaná kvalitatívna analýza existujúcich mobilných BI produktov. Výsledky analýzy, informácie o užívateľoch a kontexte použitia mobilnej BI a princípy návrhu mobilných aplikácií pre jednotlivé mobilné operačné systémy slúžili ako základ pre formuláciu pravidiel. Tie sú rozdelené do troch kategórií: stratégia doručovania BI na mobilné zariadenia, užívateľská prívetivosť a prístup k informáciám.

Abstract

This paper introduces the user-centered perspective on mobile Business Intelligence (BI). The purpose of the study was to create a set of user experience design guidelines for delivery of BI to mobile devices. First, the information about users and context of use of mobile BI were gathered based on interviews and literature review. Next, the structured qualitative analysis of fifteen current mobile BI tools was performed. The results of the study and the user experience design guidelines of four major mobile operating systems served as a basis for creating the set of guidelines for mobile BI. These guidelines are divided into three categories: Delivery strategy, Mobile user experience and Access to information.

Kľúčové slová

Mobilná business inteligencia, užívateľská prívetivosť, použiteľnosť, užívateľské rozhrania, princípy návrhu, mobilné platformy

Keywords

Mobile Business Intelligence, User Experience, Usability, User Interfaces, Design Guidelines, Mobile Platforms

1. Úvod

Mobilná business inteligencia (z angl. Business Intelligence, BI) je prostriedok, ktorý organizáciám umožňuje doručovať obsah BI reportov a dashboardov na mobilné zariadenia (telefóny a tablety) v interaktívnom, prípadne editovateľnom móde, pričom sú využité špeciálne možnosti mobilných zariadení, ako dotykové ovládanie, integrácia dát zo senzorov na rozpoznanie kontextových informácií o výpočte alebo práca v režime bez pripojenia k internetu [7]. V spojitosti s BI je *report* základný stavebný prvok dashboardu nesúci informácie, napríklad tabuľka alebo graf s dátami. *Dashboard* je interaktívny dokument zostavený z reportov a ďalších objektov, ktorý okrem dát v jednotlivých reportoch nesie vrstvu informácií vytvorených vhodnou kombináciou reportov.

Mobilita hrá v oblasti BI dôležitú úlohu. Doručovanie relevantných informácií správnym ľuďom a vo vhodnej chvíli je v komerčnej sfére kľúčové, pretože umožňuje objektívne rozhodovanie a plánovanie [3]. Vďaka prenosným zariadeniam môžu užívatelia pristupovať k dátam kedykoľvek a kdekoľvek, nezávisle od fyzickej prítomnosti na pracovisku. Na tento trend reagujú dodávatelia klasických BI produktov vývojom špecializovaných aplikácií pre mobilné platformy.

Problém je, že pre návrh takýchto aplikácií zatiaľ neexistuje rozvinutá báza znalostí v podobe návrhových vzorov alebo osvedčených postupov. Organizácie sa orientujú na implementačné otázky (Ktoré platformy podporovať? Aké technológie použiť pri vývoji?). Menej zohľadňujú potreby

užívateľov a to, ako budú aplikácie strategicky zapadať do ich pracovných procesov. Podľa prieskumov trhu [14] je pritom práve použiteľnosť a užívateľská prívetivosť mobilných BI aplikácií hlavným predpokladom pre ich adaptáciu.

Tento článok sa venuje konceptu mobilnej BI z pohľadu užívateľskej prívetivosti (voľný preklad z angl. User Experience, UX). Cieľom štúdie bolo zostavenie sady princípov pre doručovanie BI obsahu na mobilné zariadenia. Prvá časť článku sa venuje súčasnému stavu problematiky. Ďalšia časť zhŕňa poznatky o užívateľoch a kontexte použitia mobilnej BI, ktoré boli zostavené na základe prieskumu literatúry a rozhovorov s užívateľmi a produktovými dizajnérmi BI platformy. Tretia časť je venovaná štrukturovanej kvalitatívnej analýze existujúcich mobilných BI nástrojov a záverečná časť predstavuje sadu princípov doručovania BI obsahu na mobilné zariadenia.

2. Súčasný stav problematiky mobilnej BI

Návrh a vývoj mobilných BI aplikácií komplikuje niekoľko faktorov. Jednak sú to problémy spojené všeobecne s implementáciou programov pre mobilné zariadenia, napríklad fragmentácia vývoja pre rôzne platformy, nutnosť udržiavať a spravovať niekoľko aplikácií a rozdiely návrhových vzorov jednotlivých operačných systémov. Okrem toho sa tu ale vyskytujú aj komplikácie špecifické pre oblasť BI. Stipic a Bronzin vo svojom článku [12] identifikovali 4 oblasti zlepšovania mobilnej BI: hardware a komunikácia, software, architektúra a UX. Akademická sféra sa sústreďí prevažne na prvé tri okruhy, najmä bezpečnosť komunikácie [9] a architektúru [13]. UX v súvislosti s BI zostáva v pozadí napriek tomu, že v praxi patrí k najdôležitejším faktorom adaptácie medzi užívateľmi [14].

Do kategórie UX spadá jednak použiteľnosť aplikácie, teda pragmatické charakteristiky spojené s ovládaním užívateľského rozhrania (napr. naučiteľnosť, efektívnosť a zapamätateľnosť [11]). Tieto aspekty možno vyhodnocovať na základe merania rýchlosti alebo počtu chýb, ktorých sa užívateľ dopustí pri vykonávaní úlohy. Patria sem ale aj abstraktnejšie princípy týkajúce sa potrieb, motivácií a pocitov užívateľa a toho, nakoľko aplikácia spĺňa jeho očakávania a pomáha mu dosahovať stanovené ciele.

Princípy týkajúce sa návrhu UX (User Experience Design, UXD) sú súčasťou dokumentácie všetkých štyroch rozšírených mobilných operačných systémov (Android [1], iOS [8], BlackBerry [4] a Windows Phone [10]). Ich účelom je poskytnúť dizajnérovi a vývojárovi základný prehľad o konvenciách platformy a smerovať ich k tomu, aby vytvárali aplikácie konzistentné s pravidlami daného operačného systému. Dodržiavanie konvencií má priaznivý účinok na užívateľskú prívetivosť aplikácií a formuje očakávania užívateľov platformy. Pravidlá pre jednotlivé systémy a úroveň ich detailov sa však podstatne odlišujú. Kým napríklad BlackBerry uvádza len základné rozdiely mobilných aplikácií oproti desktopovým programom a niekoľko doporučení pre návrh mobilných verzií, iOS poskytuje detailný popis charakteristík platformy, vzhľadu a správania aplikácií aj jednotlivých prvkov užívateľského rozhrania. Roztrieštenosť princípov komplikuje návrh aplikácií, ktoré majú fungovať na rôznych typoch zariadení.

Ďalší problém je, že princípy návrhu aplikácií pre jednotlivé operačné systémy neposkytujú dostatočnú bázu znalostí pre návrh špecializovaných aplikácií. Nepokrývajú zásady rozdelenia funkcionality multiplatformných služieb [15] a smerujú k návrhu jednoduchých aplikácií s jedinou primárnou úlohou. Motiváciou pre vývoj jednoúčelových aplikácií je zohľadnenie podmienok použitia v mobilnom kontexte, na malej obrazovke prenosného telefónu. Tento predpoklad ale neberie do úvahy odlišné vzory používania a formáty tabletov, lepšiu podporu ovládania gestami u nových zariadení ani rastúce skúsenosti a požiadavky užívateľov.

Užívateľské rozhrania BI aplikácií sa vyznačujú vysokou informačnou hustotou a komplexnými scenármi úloh, ktoré môže užívateľ vykonávať. Medzi typické operácie patrí zobrazenie detailných informácií, porovnávanie dát, filtrovanie a drilovanie k novým informáciám, ale aj prispôbenie drilovacích ciest a filtrovacích možností a úprava samotných reportov a dashboardov. V prípade mobilnej aplikácie to znamená vysporiadať sa s tromi problémami:

- zobrazenie veľkého objemu informácií na obmedzenej ploche obrazovky prenosného zariadenia

- nízka precíznosť dotykového ovládania
- podpora pokročilých interakcií s dátami

Pre takéto úlohy na dotykových obrazovkách zatiaľ neexistujú návrhové vzory. Chýba ucelený výskum, ktorý by sa zaoberal problematikou dotykového ovládania vizualizácií a neexistujú ani knižnice prvkov, ktoré by umožnili vykresľovanie plne interaktívnych grafov na mobilných zariadeniach. Komplikáciou v tomto ohľade je aj tesné zviazanie mobilných BI aplikácií s webovými alebo desktopovými verziami. Interakcie v desktopových a webových aplikáciách sa podstatne odlišujú od tých, ktoré sú dostupné na mobilných zariadeniach. U dotykových rozhraní nemožno počítať s presným výberom položiek ako v prípade použitia klávesnice a myši, zobrazením doplnujúcich informácií pridržením kurzora nad objektom, otváraním kontextového menu ani s klávesovými skratkami. Zatiaľ neexistujú štandardy ani návody, ktoré by popisovali priamočiary prenos takýchto interakcií na mobilné platformy, alebo dokonca využitie predností mobilných rozhraní.

3. Užívatelia a kontext použitia mobilnej BI

Prvým krokom k vytvoreniu princípov doručovania BI obsahu na mobilné zariadenia bol prieskum kontextu použitia mobilných BI nástrojov. Cieľom bolo získať základné informácie o užívateľoch, ich motiváciách, obavách a potrebách spojených s používaním BI v mobilnom kontexte. Požiadavky boli formulované na základe prieskumu literatúry a dvanástich rozhovorov s užívateľmi a produktovými dizajnérmi existujúcej BI platformy.

Cieľová skupina: Primárnou cieľovou skupinou mobilných BI aplikácií sú senior manažéri a výkonní manažéri. Je to skupina užívateľov, ktorá potrebuje rýchly prístup k informáciám o spoločnosti, a u ktorej má BI podporiť proces rozhodovania. Ďalšou skupinou je stredný a nižší management a pracovníci v teréne, ktorí môžu ťažiť z proaktívnych notifikácií a prístupu k aktuálnym dátam pri operatívnom plánovaní a komunikácii so zákazníkom. [5]

Motivácia nasadenia mobilnej BI: Najčastejšou motiváciou organizácií pre nasadenie mobilnej BI je zvýšenie kvality a rýchlosti doručovania informácií a podpora rozhodovacieho procesu [5]. Ďalšia skupina motivačných faktorov je spojená s vývojom mobilného priemyslu. Prenosné zariadenia sú stále výkonnejšie, cenovo dostupnejšie a zlepšuje sa ich užívateľská prívetivosť. Vďaka flexibilitě viac vyhovujú potrebám práce mimo kancelárie a umožňujú človeku prístup k relevantným informáciám bez toho, aby bol viazaný na fyzické pracovisko. Informácie tak môžu mať k dispozícii aj pracovníci, ktorí k tradičným BI nástrojom prístup nemajú. [6]

Riziká: Na druhej strane, potenciálnych užívateľov mobilných BI aplikácií zatiaľ odrádza niekoľko rizikových faktorov. V prvom rade je to bezpečnosť informácií na prenosných zariadeniach [5]. Už zo svojej povahy sú mobilné zariadenia náchylnejšie na stratu alebo odcudzenie. Navyše, na rozdiel od telefónov, tablety bývajú často zdieľané medzi viacerými užívateľmi [2]. Zvyšuje sa tak riziko, že sa k citlivým informáciám dostane nepovolaná osoba. Komplikáciou je aj vnímanie bezpečnosti prenosu dát cez mobilné siete. Hoci existujú mechanizmy na ochranu mobilných dát [9] a ich kvalita rastie, problémy spôsobuje nízke povedomie medzi užívateľmi. Tí sa v obave pred únikom citlivých informácií vyhýbajú doručovaniu takýchto dát na prenosné zariadenia, čím znižujú hodnotu mobilného BI riešenia. Obavy sú spojené aj s použiteľnosťou a užívateľskou prívetivosťou BI nástrojov na mobilných zariadeniach. Inteligentné telefóny nemajú optimálny formát na zobrazovanie klasických BI dashboardov a manipuláciu s veľkým objemom dát [6]. Pre ovládanie vizualizácií na dotykových obrazovkách neexistujú návrhové vzory ani neformálne zaužívané praktiky. Nízka miera interaktivity u mobilných BI riešení obmedzuje pokročilých užívateľov, pretože nedovoľuje komplexné analytické operácie. V tejto súvislosti je otázkou aj miera návratnosti investície. Nasadenie a správa aplikácií a BI obsahu na rôznych platformách je pracné a dopad nie je možné priamo merať.

Požiadavky užívateľov: Podľa prieskumu trhu [5], medzi najviac žiadané funkcie mobilných BI aplikácií patrí prístup k informáciám v dashboardoch a reportoch, drilovanie, filtrovanie a proaktívne notifikácie. Pokročilé funkcie, napríklad anotácie, vkladanie dát, spúšťanie udalostí na základe zobrazených informácií, analýza dát alebo úprava reportov a dashboardov na mobilných zariadeniach,

sú v súčasnosti menej vyžadované. Záujem o ne ale stúpa s rastúcimi skúsenosťami užívateľov. Charakteristiky typické pre použitie v mobilnom kontexte, teda práca v offline režime, integrácia dát zo senzorov a prispôsobenie užívateľského rozhrania pre dotykové ovládanie, užívatelia zatiaľ explicitne nevyžadujú [14]. Ich prítomnosť však zvyšuje užívateľskú prívetivosť aplikácie.

Priorita platforiem: Medzi užívateľmi mobilnej BI je v súčasnosti najžiadanejšia podpora operačných systémov iOS a Android. Vysokú mieru adaptácie medzi koncovými užívateľmi má BlackBerry OS, ale jeho podpora postupne klesá. Ubúda aj počet užívateľov operačného systému Windows Mobile a jeho nasledovník Windows Phone zatiaľ nezískal výraznú podporu u spomínanej cieľovej skupiny. [5]

4. Analýza mobilných BI nástrojov

Druhou časťou štúdie bola štrukturovaná kvalitatívna analýza pätnástich mobilných BI riešení. Cieľom bolo vyhodnotenie stratégie rôznych dodávateľov mobilných BI produktov za účelov vytvorenia uceleného prehľadu o súčasnom stave a ponúkaných možnostiach. Aplikácie boli posudzované z pohľadu užívateľskej prívetivosti ovládania na mobilných zariadeniach a rozsahu funkcií umožňujúcich prístup k dátam. Hlavný dôraz bol kladený na silné a slabé stránky užívateľských rozhraní aplikácií a inovatívny prístup k podpore použitia v mobilnom kontexte.

4.1. Výber hodnotených aplikácií

Do tejto štúdie bolo zahrnutých pätnásť aplikácií, ktoré umožňujú užívateľovi prístup k vlastnému obsahu zobrazenému vo forme personalizovaného dashboardu. Produkty sú rozdelené do troch kategórií:

- *Mobilné aplikácie osvedčených BI dodávateľov* [7]: IBM Cognos Mobile, WebFocus Mobile Faves, MicroStrategy Mobile, Oracle Business Intelligence Mobile, SAP BusinessObjects Mobile, Tableau Mobile, Spotfire Analytics
- *Aplikácie typu SaaS*: YellowfinBI, Birst Mobile, QlikView, Domo Mobile, Bime Mobile
- *Výlučne mobilné BI aplikácie*: MyBI, SurfBI, Roambi Analytics

4.2. Kritériá hodnotenia aplikácií

Kritériá hodnotenia boli vytvorené na základe rozhovorov a prieskumu literatúry. Rozdelené sú do troch skupín:

Stratégia doručovania BI obsahu na mobilné zariadenia

- *Mobilná stratégia*: Aká je hlavná úloha mobilnej aplikácie? Ako aplikácia zapadá do BI infraštruktúry?
- *Podporované platformy*: Aké platformy dodávateľ podporuje? Ako sú riešenia implementované a aký je medzi nimi rozdiel?

Užívateľská prívetivosť

- *Dotykové ovládanie*: Ako rozhranie mobilnej aplikácie využíva prednosti dotykového ovládania, napríklad gestá a špeciálne interakcie?
- *Ovládacie prvky*: Ako sú ovládacie prvky prispôsobené dotykovému ovládaniu?
- *Naučiteľnosť*: Ako dizajn aplikácie podporuje užívateľa v procese učenia sa ovládania aplikácie?

Prístup k informáciám

- *Architektúra užívateľského rozhrania*: Akými krokmi prechádza užívateľ od spustenia aplikácie po zobrazenie konkrétneho dashboardu?
- *Typy vizualizácií*: Aké typy vizualizácií sú v mobilnej aplikácii dostupné? Ako sa vizualizácie odlišujú na rôznych platformách?
- *Interaktivita vizualizácií*: Aké operácie s vizualizáciami aplikácia umožňuje? Líši sa rozsah operácií dostupných v rôznych úrovniach detailov a na rôznych platformách?

4.3. Kritériá hodnotenia aplikácií

Pri vyhodnocovaní produktov boli použité nasledujúce materiály:

- hodnotené mobilné aplikácie (verzia pre iPad, v prípade dostupnosti verzie pre iPhone a Android)
- verejne dostupné demo projekty a vzorové dashboardy, videá, tutoriály, užívateľské príručky a iné informácie o produktoch zverejnené dodávateľmi
- diskusné fóra, recenzie užívateľov produktov u oficiálnych distribútorov (Apple AppStore, Google Play)
- nezávislé hodnotenia produktov [5], [14]

U každej aplikácie boli rámcovo testované 4 typické úlohy: prístup k dashboardom, zobrazenie detailných informácií, drilovanie a filtrovanie. To poskytlo základnú predstavu o fungovaní aplikácie a odhalilo aspekty, ktoré boli predmetom ďalšieho skúmania. Výsledky kvalitatívneho hodnotenia boli štruktúrované podľa hodnotiacich kritérií, zaznamenané do tabuľky a kódované do troch kategórií: pozitívne a negatívne príklady a netradičné riešenia.

4.4. Zhrnutie výsledkov

Z 15 hodnotených produktov slúži mobilná aplikácia v 12 prípadoch len ako doplnkový spôsob doručovania BI obsahu. U týchto riešení je na mobilné zariadenie doručovaný obsah vytváraný pre desktopové rozhranie. Cieľom je rýchla mobilizácia BI obsahu bez zvýšenia administratívnej námahy spojennej s vytváraním a prispôbovaním obsahu pre iné koncové zariadenia. Použitelnosť a užívateľská prívetivosť takýchto riešení je nízka. Obsah väčšinou nie je adekvátny pre zobrazenie na menších obrazovkách mobilných zariadení a interakcie nie sú prispôbené dotykovému ovládaniu. Z 3 zvyšných produktov spadajú 2 do kategórie čiste mobilných riešení a jeden dodávateľ poskytuje platformu pre vytváranie vlastných mobilných BI aplikácií.

Vo väčšine prípadov nie sú využité žiadne prednosti mobilných zariadení. Niektoré aplikácie ponúkajú podporu práce v offline režime, ale často sa jedná ale len o prístup k dátam, ku ktorým užívateľ pristupoval pri poslednej návšteve s pripojením k internetu.

Informácie o kontexte výpočtu využívajú 4 dodávatelia, jedná sa o geolokačné údaje a v 2 prípadoch o dáta z fotoaparátu. Podľa spôsobu implementácie možno hodnotené aplikácie rozdeliť do troch skupín: *natívne, webové a hybridné* riešenia.

Webové aplikácie umožňujú prístup k BI obsahu prostredníctvom webového prehliadača z rôznych typov zariadení. Výhodou je jednoduchá údržba a agilnejší vývoj nových verzií systému. Problém je, že webová aplikácia vyžaduje neustály prístup k internetu a prenos veľkého objemu dát, užívateľské rozhranie nie je prispôbené konvenciám a obmedzeniam koncových zariadení, reaguje pomalšie a interakcie nie sú prispôbené dotykovému ovládaniu. Len jeden z dodávateľov poskytuje webovú službu ako svoje primárne riešenie pre podporu mobilných zariadení.

Natívne aplikácie sú vytvárané samostatne pre každú podporovanú platformu. Výhodou je lepšie prispôbenie UX a širšie možnosti interaktivity. Aplikácia môže dôslednejšie rešpektovať konvencie danej platformy, výkon môže byť optimalizovaný pre daný typ koncových zariadení a využívať možno aj špecifické charakteristiky mobilných zariadení ako znalosť kontextu, alebo použitie bez pripojenia k internetu. Nevýhodou je vývoj a správa niekoľkých aplikácií a vysoké nároky na administráciu BI obsahu. Len jeden dodávateľ ponúka čisto natívnu aplikáciu prispôbenú pre konkrétnu platformu. Aplikácia patrí do kategórie čisto mobilných riešení a podporuje len zariadenia s operačným systémom iOS.

Hybridné aplikácie zobrazujú BI obsah v natívnom kontajneri. Takéto riešenie je u mobilných BI aplikácií najčastejšie. Typicky to znamená, že je navigácia aspoň čiastočne prispôbená konvenciám danej platformy a dashboardy a reporty sa vykresľujú ako webový obsah. To vedie k podstatnému obmedzeniu interaktivity. Vizualizácie nie sú prispôbené dotykovému ovládaniu alebo sú zobrazované ako statické obrázky.

Každý z dodávateľov poskytuje zákazníkovi aspoň mobilnú aplikáciu pre iPad. Druhé najpodporovanejšie zariadenie je iPhone (11 dodávateľov). Menej častá je podpora zariadení s operačným systémom Android (8 dodávateľov). Len v 3 prípadoch je dostupná aplikácia pre BlackBerry OS, vo všetkých prípadoch sa jedná o dodávateľov osvedčených BI riešení. Tablety s operačným systémom BlackBerry a zariadenia s operačným systémom Windows Phone nie sú podporované aplikáciou špeciálne vytváranou pre daný systém. Dodávatelia odkazujú zákazníkov na webové rozhranie.

Problémy s použiteľnosťou a užívateľskou prívetivosťou sú špecifické pre každú aplikáciu. Vo väčšine prípadov sa ale viažu na nešetrný prenos BI obsahu určeného pre desktop na mobilné zariadenia, nízku mieru prispôsobenia dotykovému ovládaniu a ignorovanie konvencií cieľových platforiem.

5. Princípy doručovania BI na mobilné zariadenia

Analýza existujúcich produktov slúžila ako základ pre vytvorenie sady princípov doručovania BI obsahu na mobilné zariadenia. Silné a slabé stránky jednotlivých nástrojov boli hodnotené v kontexte informácií o užívateľoch a používaní mobilnej BI a zároveň porovnávané s UXD princípmi jednotlivých mobilných operačných systémov. Na základe výsledkov hodnotenia bola zostavená iniciálna sada princípov, ktorá bola ďalej prispôbena na základe spätnej väzby od troch produktových dizajnérov BI platformy.

Princípy sú kategorizované podľa hodnotiacich kritérií analýzy. Postupujú od strategických doporučení cez princípy prispôsobenia užívateľského rozhrania dotykovému ovládaniu k špecifickým aspektom doručovania vizualizácií na mobilné zariadenia.

5.1 Stratégia doručovania BI obsahu na mobilné zariadenia

Mobilná stratégia

Prioritizácia funkcií podľa typu koncového zariadenia. Kontext použitia mobilných telefónov a tabletov sa zásadne odlišuje. Rozhranie BI aplikácie musí túto skutočnosť rešpektovať a umožňovať primárne také funkcie, ktoré sú primerané pre daný kontext. Aplikácie na mobilných telefónoch sú vhodné predovšetkým na doručovanie notifikácií a rýchly prístup k informáciám. Tablety môžu poskytovať prístup k interaktívnemu obsahu a pokročilým analytickým funkciám. Desktopové rozhranie je optimálne na vytváranie a správu BI obsahu.

Vytváranie BI obsahu. Proces vytvárania reportu sa skladá z niekoľkých krokov: formulácia cieľa, výber dát a prispôbenie ich štruktúry, výber vhodnej vizuálnej prezentácie, jej prispôbenie a nastavenie väzieb na ďalšie informácie. Správca obsahu by mal mať možnosť pracovať s vytvoreným reportom ako s objektom s tým, že miera a spôsob interaktivity by sa prispôbovala podľa toho, na akom zariadení sa report zobrazuje. Obsah a mieru interaktivity dashboardov je vhodné prispôbovať pre každú kategóriu koncových zariadení.

Sémantické informácie o dashboarde. Pri automatickom vykresľovaní obsahu dashboardu na iný typ koncového zariadenia dochádza k zmenám rozloženia informácií tak, aby zobrazenie vyhovovalo charakteristikám daného typu zariadenia. Kvôli zachovaniu informačnej hodnoty dashboardu je nutné podporiť takúto úlohu sémantickými informáciami o objektoch na dashboarde, ich poradí a vzájomných vzťahoch.

Podporované platformy

Strategický výber podporovaných platforiem. Prispôbovanie BI obsahu rôznym typom koncových zariadení je pracné a časovo náročné. Automatizácia procesu, napríklad vytvorením hybridnej aplikácie, ktorá v natívnom kontajneri zobrazuje webový obsah, nie je optimálna kvôli vysokej miere interaktivity BI aplikácií. Z hľadiska užívateľskej prívetivosti aplikácií je vhodnejší strategický výber menšieho počtu podporovaných platforiem na základe preferencií cieľovej skupiny užívateľov a vývoja trhu.

Primárny dôraz na podporu tabletov. Tablety zachovávajú výhody prenosných zariadení, ale vďaka väčšej obrazovke umožňujú prehľadnejšie zobrazenie veľkého objemu dát a podporujú zložitejšie interakcie.

5.2 Užívateľská prívetivosť

Dotykové ovládanie

Priama manipulácia s objektami. Priama manipulácia s užívateľským rozhraním je pre človeka prirodzenejšia a zvyšuje jeho pocit kontroly nad aplikáciou [1], [8]. Pokiaľ je to možné, užívateľ by mal mať možnosť používať gestá na navigáciu a ovládanie aplikácie, napríklad listovanie (gesto "swipe") na prechádzanie medzi stranami dokumentu namiesto prepínania odkazom. *Špeciálne dotykové interakcie.* Použitie špeciálnych interakcií určitej platformy, napríklad potiahnutie obsahu ("pull-to-refresh") na načítanie aktuálnych dát, podporuje konzistenciu aplikácie s pravidlami platformy.

Reakcia na zmenu orientácie zariadenia. Aplikácia by mala reagovať na zmeny orientácie zariadenia prispôbením zobrazenia na výšku a na šírku. [8] Obsah musí byť v každom móde rozmiestnený tak, aby efektívne využíval rozmery obrazovky. Ak sú v jednom prípade dostupné informácie alebo funkcie, ktoré užívateľ v inom móde nevidí, mal by byť na túto skutočnosť explicitne upozornený.

Spätná väzba prostredníctvom animácií. Vhodne použité drobné animácie poskytujú užívateľovi spätnú väzbu a umocňujú jeho pocit kontroly nad aplikáciou. [8]

Ovládacie prvky

Optimalizácia ovládacích prvkov pre dotykové ovládanie. Obzvlášť u rozhraní s veľkým množstvom interaktívnych objektov musia byť ovládacie prvky (tlačítka, odkazy a podobne) dostatočne veľké na to, aby umožnili dotykové ovládanie. Medzi interaktívnymi prvkami musí byť dostatok priestoru aby nedochádzalo k neúmyselnému výberu okolitých prvkov.

Optimalizácia textu. Veľkosť fontu a podrobnosť textu a textových označení musí byť prispôbená zvlášť pre každú kategóriu koncových zariadení, aby bola zachovaná čitateľnosť a prehľadnosť textu.

Redukcia počtu ovládacích prvkov. Počet ovládacích prvkov možno eliminovať niekoľkými spôsobmi: umiestniť na každú obrazovku len ovládacie prvky, ktoré bezprostredne súvisia s typom obsahu a úrovňou detailu; zhľukovať súvisiace prvky do skupín a sprístupniť ich na vyžiadanie; nahradiť ovládacie prvky priamou manipuláciou s obsahom; ukrývať ovládacie prvky, ktoré nie sú kriticky potrebné po určitom čase a zobraziť ich na vyžiadanie.

Vizuálne odlišenie ovládacích prvkov. Užívateľské rozhrania BI aplikácií obsahujú často veľké množstvo informácií. Ovládacie prvky musia byť jasne odlišené aby boli rozpoznateľné od samotného obsahu a zároveň nesmú byť príliš výrazné, aby potláčali dôležitý obsah do úzadia.

Naučiteľnosť

Interná a externá konzistencia. Elementy užívateľského rozhrania s rovnakým významom musia vyzeráť a správať sa rovnako v rámci celej aplikácie. Fungovanie a vzhľad ovládacích prvkov musí rešpektovať konvencie danej platformy. Dodržiavanie týchto princípov zvyšuje užívateľskú prívetivosť aplikácie, pretože umožňuje užívateľovi použiť znalosti o aplikácii a platforme. [1], [8]

Konzistencia verzií aplikácie na rôznych platformách. Vzhľad a správanie aplikácie sa môžu na rôznych platformách v detailoch odlišovať, ale vizuálna identita a princípy fungovania musia byť jednotné, aby systém pôsobil ako celok.

Nápoveda. U špecializovaných aplikácií je dôležité poskytnúť užívateľovi prístup k nápovede. Užitočným riešením je kombinácia vyhľadávania a kontextovej nápovedy [1] k obrazovke, ktorá obsahuje výlučne popis prvkov a tipy viažuce sa k danej obrazovke. *Tipy k ovládaniu aplikácie.* Tipy sú stručné správy alebo ikony, ktoré informujú užívateľa o skutočnostiach, ktoré by inak mohol prehliadnuť, napríklad použitie určitého gesta na vyvolanie udalosti. Nesmú slúžiť ako riešenie problémov s použiteľnosťou aplikácie. Ich úlohou je upozorniť užívateľa na efektívnejšie cesty použitia aplikácie.

5.3 Prístup k informáciám

Architektúra užívateľského rozhrania

Redukcia počtu úrovní v navigácii. Štruktúra dokumentov v tradičných BI nástrojoch býva často pomerne zložitá kvôli kategorizácii obsahu a prístupovým právam. Na mobilnom zariadení musí byť užívateľ od tohto oslobodený. Mal by mať prístup len k dokumentom, ktoré sú pre neho bezprostredne relevantné a potrebuje ich používať v mobilnom kontexte. Počet úrovní v štruktúre by mal byť eliminovaný vhodnou vizuálnou prezentáciou obsahu.

Vizuálne pomôcky na zjednodušenie navigácie. Použitie vizuálnych pomôcok (náhľady dashboardov, ikony, farebné kódovanie sekcií, ...) urýchľuje a zjednodušuje orientáciu v systéme. Náhľady dashboardov sú v mobilných BI aplikáciách obzvlášť užitočné kvôli dlhým názvom dokumentov.

Pokročilé možnosti vyhľadávania. Ak užívateľ pracuje v aplikácii s veľkým množstvom dokumentov, je vhodné poskytnúť mu pokročilé možnosti vyhľadávania. Sprístupniť možno podľa štruktúry obsahu napríklad hľadanie dashboardu na základe kľúčových slov, filtrovanie dokumentov podľa autora, projektu a označenia alebo radenie na základe názvu, dátumu poslednej návštevy, editácie a podobne.

Hierarchia obsahu zobrazeného na dashboarde a reporte. Dashboard má užívateľovi slúžiť na to, aby získal rýchly prehľad o aktuálnej situácii. Vizualizácie na dashboarde nemusia vždy obsahovať detailné hodnoty dát. Podľa potreby by ale užívateľ mal mať možnosť zobraziť detaily a približovať jednotlivé reporty. Zobrazenie reportu na samostatnej obrazovke zaručí lepšie využitie priestoru na uvedenie detailov a pokročilé operácie.

Typy vizualizácií

Zachovanie rovnakých typov vizualizácií na rôznych zariadeniach. Každý typ vizualizácie by mal byť dostupný na všetkých typoch koncových zariadení a to so zachovaním pokiaľ možno rovnakého rozsahu interaktivity optimalizovanej pre konkrétnu kategóriu zariadenia.

Využitie výhod dotykového ovládania. Pri vytváraní BI obsahu je nutné zohľadňovať cieľ vizualizácie a prihliadať na to, ako môže daný typ koncového zariadenia tento cieľ podporiť. Napríklad u vizualizácie trendov je častou úlohou porovnávanie hodnôt v určitých časoch. Na desktopovom rozhraní nie je takáto úloha priamočiara, pretože kurzorom užívateľ vyberá len jednu položku v danej chvíli. U dotykových rozhraní ale môžu byť efektívne využité gestá na výber viacerých položiek ("multi-finger tap" a "multi-finger drag").

Animácie zmien. Animácie môžu u vizualizácií komunikovať rozdiely medzi hodnotami. Ich použitie je u BI aplikácií na mobilných zariadeniach dôležité, pretože načítanie nového obsahu môže byť kvôli objemu sťahovaných dát pomalé. Typickým riešením u existujúcich produktov je nahradenie vizualizácie animáciou, ktorá naznačuje sťahovanie dát a následne vykreslenie kompletne novej vizualizácie. Ak sa aktualizácia dát zobrazí ako zmena oproti pôvodným hodnotám priamo vo vizualizácii, užívateľ nestratí kontext a naopak získa doplňujúce informácie o zmene.

Interaktivita vizualizácií

Zvýraznenie interaktívnych položiek v grafe. Interaktívne položky vo vizualizáciách musia byť jasne odlišné od neaktívnych prvkov. U vizualizácií určených pre desktopové rozhrania sa počíta so zvýraznením položiek pri pohybe kurzora nad interaktívnou oblasťou. Takéto odlišenie nie je u dotykových rozhraní možné.

Odstránenie redundantných interakcií. Vizualizácie nesmú obsahovať redundantné interakcie, napríklad nutnosť explicitne zatvoriť okno s detailnými informáciami o položke pred zobrazením ďalšieho detailu. Jedinou výnimkou sú operácie, u ktorých redundantná interakcia zabraňuje chybe (napr. potvrdenie kritickej operácie).

Prispôbenie interakcií nízkej precízности dotykového ovládania. Grafy a tabuľky zobrazujú veľké množstvo dát na malom priestore. Dotykové ovládanie musí byť prispôbené pre prácu s takýmto obsahom. Interakcie musia podporovať cieľ konkrétnej úlohy, napríklad pri zobrazovaní detailu o položke v grafe je namiesto stlačenia konkrétnej položky (gesto "tap") vhodnejšie dovoliť užívateľovi

prechádzať cez vizualizáciu (gesto "press and drag") a zobrazovať detaily o aktivovaných položkách. Užívateľ tak má väčšiu kontrolu nad ovládaním.

Vhodné alternatívy za interakcie prebrané z desktopových rozhraní. Tradičné desktopové a webové BI aplikácie využívajú vo veľkej miere interakcie ako zobrazenie detailov pri prechode kurzora cez interaktívnu oblasť alebo otvorenie kontextovej ponuky kliknutím na položku pravým tlačidlom myši. Za takéto interakcie musia byť u dotykových rozhraní dosadené vhodné alternatívy. Podľa konkrétnej situácie možno použiť napríklad dlhé stlačenie položky (gesto "press") alebo priblíženie položky (gesto "spread") na zobrazenie detailu, vytiahnutie ponuky (gesto "swipe") z okraja obrazovky a podobne.

6. Záver

Cieľom tejto štúdie bolo vytvorenie sady princípov doručovania BI obsahu na mobilné zariadenia. Prvým krokom bola formulácia poznatkov o užívateľoch a kontexte použitia mobilných BI nástrojov. Ďalej nasledovala analýza existujúcich mobilných BI produktov. Výsledky analýzy, informácie o užívateľoch a kontexte použitia mobilnej BI a UXD princípy jednotlivých mobilných operačných systémov slúžili ako základ pre vytvorenie prvej verzie pravidiel. Tie sú rozdelené do troch kategórií. Postupujú od strategických doporučení cez princípy prispôsobenia užívateľského rozhrania dotykovému ovládaniu k špecifickým aspektom doručovania vizualizácií na mobilné zariadenia. Táto formulácia nie je konečná. Pravidlá sa budú ďalej iteratívne vyvíjať a budú slúžiť ako základ pri vytváraní knižnice interaktívnych vizualizácií pre mobilné zariadenia.

7. Literatúra

- [1] Android Design [online] c2012 [cit. 2012-10-20] WWW: <<http://developer.android.com/design/index.html>>
- [2] Budiu, B.R., Nielsen, J.: *Usability of iPad Apps and Websites, 2nd Edition*. Nielsen Norman Group (2011).
- [3] Derballa, V., Pousttchi, K.: *Extending knowledge management to mobile workplaces*. Proceedings of ICEC '04. p. 583. ACM Press (2004).
- [4] Design principles for BlackBerry Devices [online] c2012 [cit. 2012-10-20] WWW: <http://docs.blackberry.com/en/developers/deliverables/5716/Design_principles_for_BlackBerry_devices_379328_11.jsp>
- [5] Dresner, H., et al.: *Mobile Business Intelligence Market Study*. Dresner Advisory Services (2011).
- [6] Evelson, B.: *A Practical How-To Approach To Mobile BI*. Forrester (2011).
- [7] Hagerty, J., et al.: *Magic quadrant for business intelligence platforms*. Gartner (2012).
- [8] iOS Human Interface Guidelines [online] c2010 [cit. 2012-10-20] WWW: <<http://developer.apple.com/library/ios/#documentation/UserExperience/Conceptual/MobileHIG/Introduction/Introduction.html>>
- [9] Kuntze, N., et al.: *Secure Mobile Business Information Processing*. Proceedings of 2010 IEEE/IFIP. pp. 672–678. IEEE (2010).
- [10] Microsoft: Designing UX for apps [online] c 2012 [cit. 2012-10-20] WWW: <<http://msdn.microsoft.com/en-us/library/windows/apps/hh779072>>
- [11] Nielsen, J.: *Usability 101: Introduction to Usability*. Jacob Nielsen's Alertbox (2003).
- [12] Stipic, A., Bronzin, T.: *Mobile BI: The past, the present and the future*. Proceedings of MIPRO 2011. pp. 1560–1564. IEEE (2011).
- [13] Sajjad, B., et al.: *An open source service oriented Mobile Business Intelligence Tool (MBIT)*. Proceedings of ICICT '09. pp. 235–240. IEEE (2009).
- [14] Tapadinhas, J.: *Critical Capabilities for Mobile BI*. Gartner (2012).
- [15] Wäljas, M., et al.: *Cross-platform service user experience: A Field Study and an Initial Framework*. Proceedings of MobileHCI '10.

Logování pro novou generaci monitoringu

Daniel Tovarňák, Tomáš Pitner

Masarykova univerzita, Fakulta informatiky, laboratoř Lasaris
Botanická 68a, 60200 Brno, Česká republika
{xtovarn, tomp}@fi.muni.cz

Abstrakt

Podoba a možnosti současného monitoringu začínají být ovlivňovány rostoucí variabilitou jeho využití a částečně také nástupem nových paradigmat jako je Cloud computing. Mezi výčet pozorovaných nedostatků patří například problémy s výkonností, rozšiřitelností, či interoperabilitou. Hlavním zaměřením tohoto příspěvku je problematika počítačových logů, které představují hojně využívaný typ monitorovací informace. Právě současný stav logování lze považovat za jeden z důvodů zmiňovaných nedostatků. Naším cílem je představit stávající podobu této problematiky a v návaznosti také možná řešení.

Abstract

Present-day monitoring is being affected by increasing variability of its use as well as by emergence of new computing paradigms such as Cloud computing. The list of reported shortcomings includes, but is not limited to performance, extensibility and maintainability issues. This paper is focused on computer logs – frequently used type of monitoring information. Logging is considered one of the causes of abovementioned shortcomings. Our goal is to present current work in this area as well as discuss possible solutions.

Klíčová slova

monitoring, JSON, logování

Keywords

monitoring, JSON, logging

1. Úvod

Důležitost monitoringu v poslední dekádě plynule rostla a s nástupem Cloud computingu je tento trend ještě markantnější. V rozsáhlých infrastrukturách jako je Cloud, nebo Grid je monitoring využíván nejen na provozní úrovni, ale také na úrovni služeb. Monitorovací data jsou často využívána v širokém spektru úloh, jako například accounting, analýza výkonu, audit, detekce chyb, ladění, optimalizace, či plánování. Navíc množství monitorovacích dat, které moderní systémy produkují, rapidně roste.

Míra a variabilita využití monitorovacích dat a samotné principy Cloudu začínají pomalu odhalovat nedostatky současných přístupů a technologií. Například z pohledu výkonnosti se v poslední době ukazují problémy s rychlostí zpracování monitorovacích dat, latencí a škálovatelností. Výjimkou ovšem také nejsou problémy s rozšiřitelností a interoperabilitou. V některých případech navíc současné technologie vůbec neposkytují požadovanou funkcionalitu. Toto platí především v případě podpory více uživatelů, nebo pokročilé korelace (např. techniky založené na zpracování komplexních událostí – Complex Event Processing).

V tomto příspěvku se zaměříme především na problematiku logování, jakožto jednu z podob monitorovací informace. Z mnoha důvodů (jež budou objasněny níže) lze právě počítačové logy považovat za významný faktor, jenž negativně ovlivňuje současnou podobu monitoringu. Naším cílem je diskutovat současný stav a existující přístupy. Dalším důležitým cílem je představit možné řešení a diskutovat jej v kontrastu s těmi existujícími.

2. Terminologie

Tento příspěvek se drží klasické definice monitorovacího procesu uvedeného v [1]. Proces je dělen do čtyř základních fází. Pro komponenty, které v procesu monitoringu vystupují, dále používáme terminologii původně definovanou v [2]:

2.1 Proces monitoringu

- **produkce** – získávání a tvorba monitorovacích dat
- **zpracování** – může proběhnout před nebo po kterékoliv jiné fázi procesu
- **distribuce** – přenos monitorovacích dat všem zainteresovaným stranám
- **konzumace** – finální fáze, jež zahrnuje například vyhodnocení a vizualizaci dat

2.2 Komponenty

- **senzor** – generuje prvotní monitorovací informaci
- **producent** – sbírá, zpracovává a distribuuje monitorovací data pomocí definovaného API
- **konzument** – dále zpracovává, vyhodnocuje a případně vizualizuje monitorovací data, pro konzumaci využívá API producenta
- **procesor** – je komponenta, která představuje jak konzumenta, tak producenta dat. Běžně se používá pro pokročilé zpracování monitorovací informace (např. korelace, či filtrování)

3. Problematika logování

Počítačové logy obsahují informace o aktivitě, nebo změně stavu daného systému, či aplikace. Obsahují informace o chybách, důležitá upozornění, ale často také data o průběhu business transakcí. Prakticky všechny softwarové systémy, ať již komerční, nebo open-source, logy v určité podobě používají. V drtivé většině případů představují jediný prostředek, jak tyto aplikace monitorovat, respektive získávat smysluplné informace o jejich činnosti. Lze však konstatovat, že v kontrastu s důležitostí počítačových logů pro management, ladění a analýzu softwarových systémů je současná podoba logování nedostatečná.

Typický počítačový log má podobu souboru obsahující prostý text. Každý řádek v souboru odpovídá jednomu logovacímu záznamu (což ovšem nemusí být pravidlem). Zjednodušeně se jedná o řetězec s ad-hoc strukturou. Standardy, jako například Syslog, či Apache Common Log Format, jež podobu tohoto řetězce do jisté míry určují, se nejenže mezi sebou netriviálně liší, ale ani jejich použití není zárukou logické struktury logovacího záznamu. Je to mimo jiné důsledkem faktu, že informace s nejvyšší informační hodnotou je zaznamenána v podobě přirozeného jazyka („*server argo.fi.muni.cz dostal požadavek z adresy 10.1.18.254*“). Rozšiřitelnost takovéto struktury je navíc minimální. Podstatnou komplikací je nakonec fakt, že logy jsou zpravidla rozmístěny napříč adresářovou strukturou a je tak třeba k nim přistupovat ad-hoc.

Vzhledem k obrovskému množství logovacích záznamů a k frekvenci s jakou v moderních systémech vznikají, je pro člověka nemožné provádět jejich manuální analýzu [3][4]. V důsledku výše uvedených faktů ovšem nelze logy přímo zpracovávat ani zcela automaticky.

Běžným uživatelským přístupem k řešení uvedených problémů je manuální tvorba proprietárních skriptů s komplexními regulárními výrazy, což s sebou nese jisté netriviální úsilí. Naopak udržitelnost a rozšiřitelnost takového řešení je velmi nízká.

Komplexnější přístupy a publikace lze rozdělit do dvou základních skupin. První skupina přístupů se zaměřuje na analýzu a zpracování již vygenerovaných logů za použití technik a algoritmů analýzy dat (především z oblasti data-miningu), tj. na straně konzumenta. Druhá skupina zahrnuje přístupy a studie, jež se zaměřují na kvalitu a prvotní formu monitorovacích dat ještě před jejich vznikem, tj. na straně producenta.

3.1 Přístupy zaměřené na analýzu počítačových logů

Základním cílem přístupů zaměřených na zpracování a analýzu logovacích záznamů s ad-hoc strukturou je odvození, nebo také *abstrakce* typů těchto záznamů. Jinak řečeno se jedná o separaci

statických (společných) rysů logů od dynamicky se měnících rysů (proměnných) [5]. Z technického pohledu je prvním cílem automatické vygenerování regulárních výrazů, resp. vzorů, jež jsou používány k dalšímu zpracování.

Prvním možným přístupem k abstrakci typů je použití technik data-miningu k odhalení frekventovaných vzorů. Kupříkladu Vaarandi [6][7] a Makanju [8] používají shlukovou analýzu (clustering) – v jednoduchosti se jedná o rozdělení objektů do skupin (shluků) tak, aby si prvky ve stejné skupině byly co nejpodobnější. Rozšířením práce Vaarandiho je [5], kdy se při shlukování bere v potaz také frekvence výskytů slov v jednotlivých záznamech.

Druhým základním přístupem je abstrakce typů s využitím analýzy zdrojových kódů, často používaný v oblasti detekce anomálií a chyb analyzovaného programu [9][10][11]. Pro úplnost uvedme hybridní přístupy, které kombinují shlukovou analýzu s analýzou zdrojového kódu [12][13].

Lze konstatovat, že přístupy zaměřené na abstrakci nejsou z principu příliš vhodné pro zpracování obrovského množství dat v reálném čase, což je z pohledu použití v oblasti Cloud computingu nevýhoda. Navíc zejména v případě shlukové analýzy se jedná o aproximaci, tzn. lze narazit na okrajové případy, které vyžadují lidskou intervenci.

3.2 Přístupy zaměřené na podobu počítačových logů

Z prací, které se zaměřují především na logování jako takové z pohledu vývoje a základních mechanismů, je třeba nejdříve zmínit [14]. Yuan zde podrobně představuje důležité charakteristiky, jež se týkají způsobů logování a podoby vlastních logovacích záznamů. Dodejme, že tyto charakteristiky byly získány analýzou čtyř velkých open-source projektů. Podobně [15] se zaměřuje na problematiku evoluce, verzování a dokumentace logovacích záznamů v průběhu vývoje. Konečně [16] se zaměřuje na bezpečnost logů z pohledu zabezpečení, integrity a soukromí. Zdůrazňuje, že důležitost těchto faktorů je akcentovaná nástupem sdílených prostředí, jako je Cloud.

4. Požadavky na producenta monitorovacích dat

V [17] byly identifikovány možné příčiny existujících problémů v monitoringu a především definovali požadavky na producenta monitorovacích dat se zaměřením na virtualizované prostředí. Na rozdíl od přístupů zaměřených na abstrakci, agregaci a konverzi monitorovací informace se tak zaměřujeme na prvotní kvalitu, podobu a způsob poskytování produkovaných dat. V následujících odstavcích krátce shrneme hlavní požadavky na producenta monitorovací informace a také uvedeme některé navrhované způsoby jak těmto požadavkům vyhovět.

4.1 Jednotná reprezentace monitorovací informace

Monitorovací data se v moderních systémech vyskytují především ve dvou hlavních podobách, a to v podobě měření a počítačových logů. Měření vyjadřují metriku sledovaného systému (zátěž CPU, volná paměť, počet běžících procesů atd.) zatímco logy nesou informaci o aktivitě či stavu daného systému. Hlavním problémem je především odlišná reprezentace těchto dat, z čehož plyne především omezená možnost jejich korelace, tj. například dotazy typu „*Jaká událost způsobila nárůst zátěže CPU o 60% v daném intervalu?*“. Možným řešením je reprezentovat všechny druhy monitorovacích dat jako události v čase. U měření to například znamená reprezentovat jednotlivé odečty jako události v čase, respektive reagovat na předem definované změny dané veličiny (např. *za poslední vteřinu vzrostla zátěž CPU o 4%*). Naopak logy samotné lze považovat za události v čase.

4.2 Strukturovaný datový formát

S předchozím bodem nedílně souvisí také formát, v jakém jsou jednotlivé události reprezentovány. Velké množství nedostatků uvedených v úvodu souvisí právě s formátem monitorovacích dat. Jak jsme již zmínili, velmi problematický je především formát počítačových logů. Nejenže jsou logovací formáty nestrukturované, ale navíc se od sebe liší také podle standardu, který implementují (např. Apache Common Log Format versus Syslog). Největším problémem je ovšem fakt, že nejdůležitější informace mají v drtivé většině případů podobu řetězce v přirozeném jazyce („*server argo.fi.muni.cz*

byl restartován“). Informaci v takovéto podobě je velmi obtížné efektivním způsobem zpracovat, což výrazně prodlužuje čas potřebný k vyhodnocení takové informace. Možnému řešení se budeme krátce věnovat dále.

4.3 Standardní protokol pro distribuci

Za předpokladu, že jsou všechny podoby monitorovací informace reprezentovány „standardním“ způsobem, je přirozené požadovat, aby byly také distribuovány jednotným způsobem. Takového stavu je možné dosáhnout například zavedením standardního protokolu pro přenos monitorovacích dat.

Z pohledu transferu dat je žádoucí, aby bylo pokryto co nejširší portfolio interakcí mezi producentem a konzumentem monitorovacích dat. V praxi to znamená především podporu obou základních komunikačních modelů – *pull*, kdy je komunikace iniciována konzumentem, a *push*, kdy je komunikace naopak iniciována producentem. Dalším požadavkem je též podpora jak *synchronního* (zaručeného), tak i *asynchronního* přenosu dat.

4.4 Podpora více uživatelů (konzumentů)

V prostředí Cloudu je běžné, že uživatelé provozují kritické aplikace na infrastruktuře, kterou nemohou plně ovládat a spravovat. Ovšem jak uživatel, tak poskytovatel infrastruktury mají enormní zájem získávat co nejširší množinu monitorovacích dat. Někdy je naopak nutné konkrétním uživatelům odepřít přístup k citlivým monitorovacím informacím. Je tedy nutné umožnit více uživatelům přistupovat souběžně ke stejným monitorovacím datům a zároveň poskytnout mechanismus pro řízení přístupu k citlivým datům. Toto lze zajistit splněním čtyř základních předpokladů.

- **souběžnost** – více konzumentů může souběžně přistupovat k totožné monitorovací informaci
- **izolace** – konzumenti mohou přistupovat pouze k monitorovacím datům, která jsou jim určena
- **integrita** – nikdo nesmí být schopen smazat, či pozměnit senzorem vygenerovanou monitorovací informaci
- **nepopiratelnost původu** – původce (senzor) konkrétní monitorovací informace musí být nepopiratelný

5. Návrh řešení

V [17] byly definovány základní požadavky na producenta monitorovacích dat. Z pohledu logování byl nejdůležitějším požadavkem strukturovaný datový formát pro reprezentaci monitorovacích dat, tedy i logovacích záznamů. Vhodný formát by měl splňovat následující podmínky:

- standardizace
- strukturovanost
- samo-popisnost
- rozšiřitelné schéma
- kompaktnost

Za vhodné kandidáty lze považovat například formáty JSON a XML, kde ale v případě XML požadavek na kompaktnost není zcela uspokojen. V porovnání s formátem JSON má stejná informace v XML až o 50% větší velikost. Pro JSON hovoří také menší výpočetní náročnost při serializaci a deserializaci. V prototypu představeném v [17] jsou monitorovací data reprezentována právě v JSON formátu. Každý monitorovací záznam (ať již log, měření, nebo jakýkoliv jiný typ monitorovací informace) má podobu JSON objektu s množinou pevných atributů (tzv. fixní schéma) a dále volitelné schéma, jež může obsahovat libovolné atributy. Příklad takového monitorovacího záznamu (události) je uveden níže.

```
{ "Event": {  
  "id": 1605,  
  "occurrenceTime": "2012-10-12T03:44:52.713+0000",  
  "detectionTime": "2012-10-12T03:44:52.791+0000",
```

```

"hostname": "domain.localhost.cz",
"type": "org.linux.cron.Started",
"application": "Cron",
"process": "cron",
"processId": 4219,
"severity": 5,
"priority": 4,
"payload": [
  {
    "http://cron.org/1.0/events.jsch": {
      "value1": 4648,
      "value2": "3df23c7"
    }
  }
]
}
}

```

Příklad. 29. Monitorovací data v podobě JSON události

Vzhledem k možnostem existujících JSON parserů, lze monitorovací data v této podobě zpracovávat rychle a efektivně. K jednotlivým atributům se přistupuje přímo, bez použití složitých regulárních výrazů. Díky použití verzovaného schématu rozšiřitelnost a interoperabilita takového řešení drasticky roste. Pro formát JSON navíc existují binární formáty (SMILE, BSON), které dále zvyšují rychlost serializace/deserializace.

Použití standardizovaného a strukturovaného formátu s sebou také nese praktické dopady na požadavky z [17], rekapitulované v sekci 4. Lze konstatovat, že formát JSON je vhodný k reprezentaci logů, měření a ostatních možných typů monitorovacích dat a vede tak na jejich jednotnou reprezentaci v podobě událostí. Díky standardizaci a kompaktnosti je tento formát vhodný pro přenos mezi producentem a konzumenty.

6. Závěr

V tomto příspěvku jsme se zaměřili na logování, jakožto nedílnou součást moderního monitoringu. Představili jsme si terminologii všeobecně používanou v kontextu monitoringu a také naše předchozí práce, z nichž čerpáme především definici čtyř základních požadavků na producenta monitorovacích dat – jednotná reprezentace, strukturovaný formát, standardní distribuční kanál a podpora více konzumentů. Představili jsme a diskutovali problematiku logování v moderních softwarových systémech včetně krátké rešerše současných přístupů a řešení.

Jako navrhované řešení bylo představeno použití formátu JSON jakožto vhodné alternativy za v současnosti používané nestrukturované řetězce. Použití vhodného formátu s sebou přináší nejen zrychlení zpracování monitorovacích dat, ale také přímo umožňuje splnění již zmíněných požadavků na jejich producenta.

Strukturovanost a použití definovaného schématu mají za důsledek výslednou jednoduchost a efektivitu zpracování jednotlivých monitorovacích záznamů. Díky tomu se otevírají nové možnosti výzkumu především v oblastech pokročilého zpracování velkých objemů dat. Z našeho pohledu jsou zajímavé především technologie na zpracování komplexních událostí, jež umožňují definici složitých vzorů pro detekci kauzálních a temporálních závislostí.

7. Literatura

- [1] M. Mansouri-Samani. Monitoring of Distributed Systems. University of London, 1995.
- [2] B Tierney, R Aydt, D Gunter, W Smith, and M Swany. A grid monitoring architecture. 2002.

- [3] W. Jiang and et al. Understanding customer problem troubleshooting from storage system logs. In Proceedings of USENIX FAST'09, 2009.
- [4] Oliner and J. Stearley. What supercomputers say: A study of five system logs. In Proc. IEEE DSN, Washington, DC, 2007.
- [5] M. Nagappan and M.A. Vouk. Abstracting log lines to log event types for mining software system logs. In Mining Software Repositories (MSR), 2010 7th IEEE Working Conference on, May 2010.
- [6] R. Vaarandi. A Data Clustering Algorithm for Mining Patterns from Event Logs. In Proceedings of the 2003 IEEE Workshop on IP Operations and Management (IPOM), 2003.
- [7] R. Vaarandi. Mining Event Logs with SLCT and Loghound. In Proceedings of the 2008 IEEE/IFIP Network Operations and Management Symposium, April 2008.
- [8] Makanju, A. N. Zincir-Heywood, and E. E. Milios. Clustering Event Logs Using Iterative Partitioning. In Proceedings of the 15th ACM Conference on Knowledge Discovery in Data., July 2009.
- [9] W. Xu, L. Huang, A. Fox, D. Patterson, and M. Jordan. Mining Console Logs for Large-Scale System Problem Detection. SysML'08, December 2008.
- [10] W. Xu, L. Huang, A. Fox, D. Patterson, and M. Jordan. Detecting large-scale system problems by mining console logs. Proceedings of the ACM SIGOPS 22nd symposium on Operating systems principles, 2009
- [11] D. Yuan, H. Mai, W. Xiong, L. Tan, Y. Zhou, and S. Pasupa- thy. SherLog: Error diagnosis by connecting clues from run- time logs. In Proceedings of the fifteenth edition of ASPLOS on Architectural support for programming languages and operating systems (ASPLOS), 2010.
- [12] Z.M. Jiang, A.E. Hassan, G. Hamann, P. Flora. Abstracting Execution Logs to Execution Events for Enterprise Applications. Journal of Software Maintenance and Evolution: Research and Practice. Volume 20 Issue 4, 2008.
- [13] Cheng Zhang, Zhenyu Guo, Ming Wu, Longwen Lu, Yu Fan, Jianjun Zhao, and Zheng Zhang. 2011. AutoLog: facing log redundancy and insufficiency. In Proceedings of the Second Asia-Pacific Workshop on Systems (APSys '11).
- [14] Ding Yuan; Soyeon Park; Yuanyuan Zhou. Characterizing logging practices in open-source software, 34th International Conference on Software Engineering (ICSE 2012), June 2012.
- [15] Weiyi Shang. Bridging the divide between software developers and operators using logs, 34th International Conference on Software Engineering (ICSE 2012), June 2012.
- [16] Ryan K.L. Ko, BuSung Lee, and Siani Pearson. Towards achieving accountability, auditability and trust in cloud computing. In Advances in Computing and Communications, volume 193 of Communications in Computer and Information Science. Springer Berlin Heidelberg, 2011.
- [17] D. Tovarňák and T. Pitner. Towards Multi-Tenant and Interoperable Monitoring of Virtual Machines in Cloud, 12th International Symposium on Symbolic and Numeric Algorithms for Scientific Computing (SYNASC 2012), Workshop on Management of Resources and Services in Cloud and Sky Computing, 2012 [pre-print]

Computation of diffusion coefficients for lipids in polymers

Jaroslav Urbánek¹, Tatsiana P. Rusina¹, Foppe Smedes^{1, 2}

¹ Masaryk University, Research Centre for Toxic Compounds in the Environment,
Kamenice 126/3, 62500 Brno, Czech Republic
{urbanek,rusina}@recetox.muni.cz

² Deltares, P.O. Box 85467, 3508 AL Utrecht, The Netherlands

Abstract

Diffusion is an important parameter for evaluation of various uptake processes in environment including application of passive sampling technique which involves use of polymers. Diffusion coefficient is a crucial parameter for estimation of permeability of polymers towards various chemicals. This value can be computed using Fick's laws of diffusion and a mathematical theory. Since every experiment is influenced by a measurement error, it is important to analyse the uncertainty of the result. This can be done for example by the Monte Carlo method.

Abstrakt

Difuze je významným parametrem při hodnocení různých procesů příjmu látek v prostředí, včetně aplikace technik pasivního vzorkování, které vyžadují použití polymerů. Koeficient difuze je zásadní parametr pro odhad prostupnosti polymerů pro různé chemické látky. Tuto hodnotu můžeme vypočítat pomocí Fickových zákonů difuze a matematické teorie. Jelikož je každý experiment ovlivněn chybou měření, je nutné analyzovat neurčitost výsledku. To lze provést například metodou Monte Carlo.

Keywords

Diffusion coefficient, diffusion equation, uncertainty, Monte Carlo method

Klíčová slova

Koeficient difuze, rovnice difuze, neurčitost, metoda Monte Carlo

1. Introduction

Passive sampling is applied in the last 20 years for monitoring of trace levels of environmental contaminants and has become a powerful tool. It involves the exposure of an organic polymer to water, sediment, soil, air or biota tissue. During continued exposure, the analyte concentration increases in the passive sampler with time until equilibrium is obtained.

For application of passive sampling in biota tissue the information on diffusion coefficients of lipids in polymers is needed for proper method development. The method will be applied for biomonitoring which is further relevant for risk assessment of priority and emerging pollutants in environment.

The present study aimed measurement of diffusion of lipids (if any) inside the selected polymer, i.e. low density polyethylene to find out if it affects diffusion and partitioning of hydrophobic environmental pollutants (PCBs, PAHs, OCPs) what will further help evaluation of the method.

The process of diffusion is modeled by the Fick's second law [1]

$$\frac{\partial C(x,t)}{\partial t} = D \cdot \frac{\partial^2 C(x,t)}{\partial x^2} \quad (1)$$

where D is the diffusion coefficient inside the polymer, which is constant (independent of time and place), $C(x,t)$ is the concentration of diffusing substance at distance x from the reference point, and t is the diffusion time.

The equation **Chyba! Nenalezen zdroj odkazů.** was solved analytically by the means of separation of variables. The diffusion coefficients were estimated by minimizing the residual sum of squares of measured concentrations and concentrations predicted by **Chyba! Nenalezen zdroj odkazů.** Further,

the uncertainty analysis was executed through the Monte Carlo method to find out how the diffusion coefficients can differ due to measurement inaccuracy.

2. Analytical solution of the diffusion coefficient

2.1 General solution of the diffusion equation

The separation of variables method assumes that the solution $C(x, t)$ of the equation **Chyba! Nenalezen zdroj odkazů.** can be written by parts as

$$C(x, t) = X(x) \cdot T(t) \quad (2)$$

where $X(x)$ depends only on variable x and $T(t)$ depends analogically only on variable t . By applying **Chyba! Nenalezen zdroj odkazů.** to **Chyba! Nenalezen zdroj odkazů.** and separating the variables we get:

$$\frac{1}{D \cdot T(t)} \cdot \frac{\partial T(t)}{\partial t} = \frac{1}{X(x)} \cdot \frac{\partial^2 X(x)}{\partial x^2} \quad (3)$$

The only case **Chyba! Nenalezen zdroj odkazů.** can hold is when both sides of **Chyba! Nenalezen zdroj odkazů.** are equal to some constant. Let us denote such constant as λ^2 . Solving the left hand side of **Chyba! Nenalezen zdroj odkazů.** equal to λ^2 leads to a solution (e.g. see [2]):

$$T(t) = A \cdot e^{D \cdot \lambda^2 \cdot t} \quad (4)$$

where A is a real constant. Solving the right hand side of **Chyba! Nenalezen zdroj odkazů.** equal to λ^2 leads to a solution (e.g. see [2]):

$$X(x) = B \cdot e^{\lambda \cdot x} + G \cdot e^{-\lambda \cdot x} \quad (5)$$

where B, G are some real constants as well. The general solution of **Chyba! Nenalezen zdroj odkazů.** is then:

$$C(x, t) = (B \cdot e^{\lambda \cdot x} + G \cdot e^{-\lambda \cdot x}) \cdot A \cdot e^{D \cdot \lambda^2 \cdot t} \quad (6)$$

which can be further simplified to:

$$C(x, t) = e^{D \cdot \lambda^2 \cdot t} \cdot (E \cdot e^{\lambda \cdot x} + F \cdot e^{-\lambda \cdot x}) \quad (7)$$

for new real constants E, F .

2.2 Particular solution of the diffusion equation

The experiment was conducted by stacking six polymer sheets together. At the beginning (at time $t = 0$ s) the top sheet was contaminated with lipid.

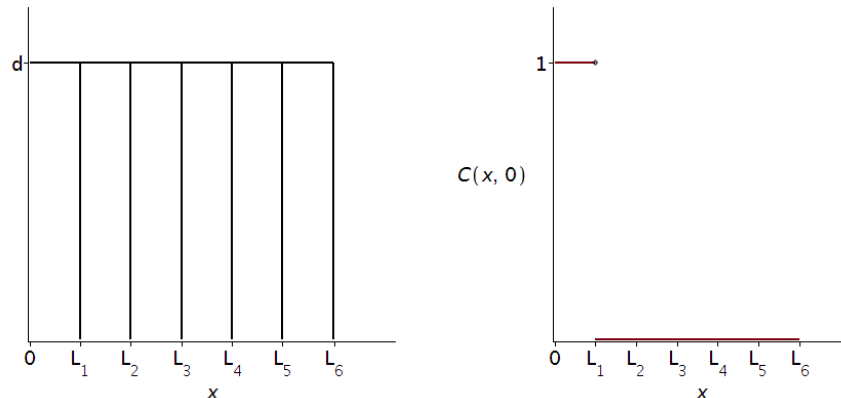


Fig. 1. *Left:* Sheets with diameter equal to d stacked together. The lengths ($L_1, L_2 - L_6, \dots$) of the sheets were similar but not precisely equal. *Right:* At time $t = 0$ s only the top sheet contained lipid.

The illustration of the experiment can be seen on the Figure 2. The sheets were surrounded by impermeable walls, which means that the chemical could move only within those six polymer sheets. For two impermeable walls the boundary conditions are specified as follows:

$$\left. \frac{\partial C(x, t)}{\partial x} \right|_{x=0} = 0 = \left. \frac{\partial C(x, t)}{\partial x} \right|_{x=L} \quad (8)$$

where L is the total length of all sheets together. Since according to **Chyba! Nenalezen zdroj odkazů.**:

$$\frac{\partial C(x, t)}{\partial x} = e^{D \cdot \lambda^2 \cdot t} \cdot (E \cdot \lambda \cdot e^{\lambda x} - F \cdot \lambda \cdot e^{-\lambda x}), \quad (9)$$

the only case **Chyba! Nenalezen zdroj odkazů.** can hold for non-primitive solution ($\lambda \neq 0$) is when λ^2 is negative. Let us therefore denote $\xi^2 = -\lambda^2$. Then the general solution **Chyba! Nenalezen zdroj odkazů.** is of the shape:

$$C(x, t) = e^{-D \cdot \xi^2 \cdot t} \cdot (A \cdot \cos(x \cdot \xi) + B \cdot \sin(x \cdot \xi)) \quad (10)$$

for some new real constants A, B . We can see (from **Chyba! Nenalezen zdroj odkazů.**) that $B = 0$, however the presence of cosine in **Chyba! Nenalezen zdroj odkazů.** leads to infinitely many solutions which can be written as:

$$C(x, t) = c_0 + \sum_{n=1}^{\infty} e^{-D \cdot \left(\frac{n \cdot \pi}{L}\right)^2 \cdot t} \cdot c_n \cdot \cos\left(\frac{n \cdot \pi \cdot x}{L}\right) \quad (11)$$

where c_n are real constants (for $n = 0, 1, 2, \dots$) which can be computed as follows:

$$c_0 = \frac{1}{L} \cdot \int_0^L C(x, 0) dx, \quad (12)$$

$$c_n = \frac{2}{L} \cdot \int_0^L \cos\left(\frac{n \cdot \pi \cdot x}{L}\right) \cdot C(x, 0) dx. \quad (13)$$

For more details about this solution, please, see e.g. [3].

2.3 Finding the diffusion coefficient value

After a given exposure time $t = t_{exp}$ the concentration $C_m(x_i, t_{exp})$ of the substance was measured in the i -th sheet for $i = 1, 2, \dots, 6$. These values express the mean concentrations in the sheets. The mean value $C_c(x_i, t_{exp})$ of the computed concentration $C(x, t)$ in the i -th sheet is given by (let $L_0 = 0$):

$$\frac{1}{L_i - L_{i-1}} \cdot \int_{L_{i-1}}^{L_i} C(x, t_{exp}) dx. \quad (14)$$

The diffusion coefficient is estimated by minimizing the following sum of squares:

$$\sum_{i=1}^6 (C_m(x_i, t_{exp}) - C_c(x_i, t_{exp}))^2. \quad (15)$$

2.4 Uncertainty analysis using Monte Carlo method

The Monte Carlo method is an algorithm based on “many” random evaluations of the model [4]. We used this method to estimate how the measurement error of concentrations $C_m(x_i, t_{exp})$ affects the obtained value of the diffusion coefficient D . Since only one measurement was made for a given

experiment, the measurement error had to be based on an expert judgment. This resulted in representation of $C_m(x_i, t_{exp})$ as a random variable with a uniform probability distribution.

Let us denote f the optimization function, which computes the minimum of **Chyba! Nenalezen zdroj odkazů..** Let $\mathbf{C}_k = (C_m(x_1, t_{exp}), \dots, C_m(x_6, t_{exp}))_k$ be the k -th realisation of random variables representing the measured concentrations $C_m(x_i, t_{exp})$ for $i = 1, 2, \dots, 6$. Then the final value of the diffusion coefficient is taken as a mean value of all model evaluations (created by Monte Carlo), which means

$$D = \frac{1}{N} \cdot \sum_{k=1}^N f(\mathbf{C}_k) \quad (16)$$

where N is the number of model evaluations. Uncertainty bounds can be reported for example in terms of standard deviation. Suitability of N can be assessed with a Monte Carlo error estimate MC_{error} given by (for more details, see [5]):

$$MC_{error} = \sqrt{\frac{1}{N \cdot (N-1)} \cdot \sum_{k=1}^N (f(\mathbf{C}_k) - D)^2} \quad (17)$$

3. Results

All computations were made in computer algebra system Maple capable of solving given problems analytically. Input data were loaded from MS Excel file for which Maple offers a package called *ExcelTools* since the version 15. Generating realisations from the uniform probability distribution was handled with the *Statistics* package and minimizing the sum of squares in **Chyba! Nenalezen zdroj odkazů..** was done by the command *Minimize* from the *Optimization* package.

We present the results obtained for the experiment with various oils (coded as “ol-3”, “ol-2”, ..., “ol+3”) diffusing inside low-density polyethylene. Main computed values are given in Table 1. The number of model evaluations N was equal to 5000, SD stands for Standard Deviation.

Table 7: Results for oils diffusing inside low-density polyethylene.

<i>chemical</i>	$D \cdot 10^{-11} (cm^2/s)$	$SD(D) \cdot 10^{-11} (cm^2/s)$	$MC_{error} \cdot 10^{-11}$
ol-3	9,44	1,60	0,023
ol-2	2,29	0,42	0,006
ol-1	1,12	0,27	0,004
ol	3,23	0,51	0,007
ol+1	2,85	0,47	0,007
ol+2	1,87	0,34	0,005
ol+3	2,29	0,41	0,006

The results from table 1 can be illustrated by plotting the function $C(x, t_{exp})$.

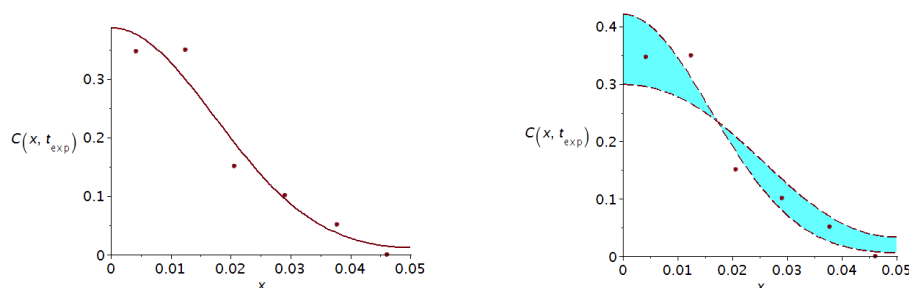


Fig. 2. The measured concentrations plotted as dots (and placed in the midpoints of intervals representing the sheet lengths – see figure 1) with computed concentration function after a certain exposure time for lipid ol-3 on the left. On the right hand side of the figure additionally the uncertainty bounds are shown.

4. Conclusion

The computed diffusion coefficients for lipids were in order of 10^{-11} being only factor 35 lower than diffusion coefficients of higher hydrophobic contaminants (PCBs/PAHs) for LDPE [6]. By the means of Monte Carlo method we can easily evaluate the uncertainties in the obtained solution. Nevertheless, the weakest point in application of the method was the determination of the uncertainty of diffusion experiments since there were not enough measurements. As Monte Carlo method relies on the experimental uncertainty including analytical error, in further work care should be taken on the proper estimation of such uncertainty by including more duplicate experiments.

Acknowledgments. This research has been supported by the CETOCOEN project from the European Regional Development Fund (No.CZ.1.05/2.1.00/01.0001).

5. References

- [1] Crank, J.: The Mathematics of Diffusion, 1st ed.; University Press: Oxford, 1957
- [2] Kalas, J., Rab, M.: Ordinary differential equations (in Czech). Masaryk University. Brno (2001)
- [3] Wolfram: Heat Conduction Equation. Web. 23 October 2012, <<http://mathworld.wolfram.com/HeatConductionEquation.html>>
- [4] Wikipedia: Monte Carlo method. Web. 23 October 2012, <http://en.wikipedia.org/wiki/Monte_Carlo_method>
- [5] Komarek, A.: Bayesian methods (in Czech). Online. 24 October 2012, <http://www.karlin.mff.cuni.cz/~komarek/vyuka/2010_11/nstp021/NSTP021-Bayes-Komarek-tisk.pdf>
- [6] Rusina, T. P., Smedes, F., Klanova, J.: Diffusion coefficients of polychlorinated biphenyls and polycyclic aromatic hydrocarbons in polydimethylsiloxane and low-density polyethylene polymers. *J. Appl. Polym. Sci.* 2010, 116 (3):1803-1810



9. letní škola aplikované informatiky
Sborník příspěvků
Bedřichov, 3.–5. září 2012

Editoři:

prof. RNDr. Jiří Hřebíček, CSc.
ing. Jan Ministr, Ph.D.
doc. RNDr. Tomáš Pitner, Ph.D.

Vydal: **Tiskárna Knopp**
Černčice 24
549 01 Nové Město nad Metují
[http:// www.tiskarnaknopp.cz](http://www.tiskarnaknopp.cz)

Tisk: **Tiskárna Knopp**
Černčice 24
549 01 Nové Město nad Metují
[http:// www.tiskarnaknopp.cz](http://www.tiskarnaknopp.cz)

ISBN