

On Creating Complex Modelling Systems Of Caves

Peter Tarsoly, PhD.

Alba Regia Technical Faculty of Óbuda University, Székesfehérvár, Hungary
tarsoly.peter@amk.uni-obuda.hu

Abstract - Caves are such underground objects, which are important not only for speleologists, but for medical treatments, security protection, or even for casual tourists as well. Basically, the knowledge on the geometry of a cave may be essential: based on that several cave-arms can be discovered, which may be interesting for several aspects. The mapping of caves is done by surveying tools and method. In practice, these are outdated techniques, which does not provide the appropriate accuracy demand of the mapping. Therefore, in the present study the viability of more recent surveying instrumentation is investigated. Subsequently, GIS database has been developed based on surveying measurements of several Hungarian caves and additional data collection related to the certain caves. The developed GIS has become a complex modelling systems of caves, which can be a valuable tool for several applications ranging from healthcare to family excursion.

Keywords: Caves, GIS, Surveying, GNSS

I. INTRODUCTION

Caves are such underground objects, which are isolated from the outer world since the beginning of the prehistoric ages. The dust-free and partly weather-independent environment allows the formation of a very special flora and fauna. The geologist can look into the secrets of the Earth by reconstructing the history of the crust. From a different aspect, caves are important for tourists; many people are amazed by the world of the eternal night. The changes on the Earth surface and the water pollution affect the sensitive balance of underground world. That is, why it is so important to create complex modelling systems of caves, which integrates spatial processes in a cave, and their links to each other.

The geographic information system (GIS) sciences have gone through dynamic changes in the last few years in Hungary. People has recognized that the protection of natural environmental sights can be achieved only by close cooperation of the different civil services. This can be supported by GIS by building a complex information system, which provides a uniform base for all technological and economical planning. Not only the national parks and nature conservation areas are claiming up-to-date digital maps, but these may also be relevant for local governmental bodies, technological and economical firms. In fact, most Hungarian caves are not secure against humans, the bluntness and fear of some common men are often destroying them – this moment is maybe the last, when we can do something to rescue them. The primer task is to build a cadastral information

system based on measured three-dimensional geometric data of the caves and the mountains for identification and registration. Subsequently, basic processes and properties of the caves can be attached to these data, which gives a description of its spatial distribution and the actual environmental effects. The emphasis is on the spatial data, since a good decision can be based only on knowing the coherence between the underground world and the surface of the Earth.

In this study I was build two cadastral information system: one about the caves of the Keszthely-mountain, and the other about the caves of the Tihanyi-peninsula. The most caves in the Keszthely-mountain are limestone/dolomite caves and basalt-caves [1, 2], and all the caves of the Tihanyi-peninsula are non-karstic caves, they are about gejsirit [2]. To create a map of a karst-area, or a non-karst area, or to measure a karst-or a non-karst cave are essentially different tasks, making use of different methods, which can be supported by an information system equipped with variable kind of attribute data. The whole workflow I can partition to four steps [4]:

1. Determination of cave entrances
2. Surveying the underground world
3. Making a map about the measured data
4. Building the complex information system

II. DETERMINATION OF CAVE ENTRANCES

Coordinates of the entrances of most caves can be determined with the Global Navigation Satellite System (GNSS). The entrances can be identify only with sub-meter accuracy, so use of GNSS receivers must be sufficiently adequate. In the summer of 2001, I was investigating the optimal measuring method in the area of the National Park of the Balatonfelvidék [3]. I was looking for the optimal measuring time and processing method in many modes of observation. Approximately, I have measured 300 cave entrances with four GPS-receivers: the Leica-500 (a receiver with geodetic precision), the Trimble ProXr (suitable for GIS measurements), and the Magellan Tracker and Garmin Etrex (navigation devices). Auxiliary instruments were a tape, a compass and a Leica laser distance-meter. Often we cannot set up the instrument right at the entrance of the cave, but in a few meter distance apart from it (Fig.1).

For measuring the direction between the instrument and the entrance we have used a compass, and for measuring the distance, a tape or a distance-meter has been used. Using the latter was difficult in strong Sunshine, since we could not identify exactly the place of

the laser-point. Even though we have measured slope distances, it was not necessary to reduce them to horizontal distances due to the short distances and the nearly horizontal measurements. Due to the masking of the trees or the hang over rocks, in most cases only the signs of four satellites could be observed, so we had to extend the measuring time from the originally estimated 10 minutes to 20-25 minutes. In the processing manner of the measurements several options have been investigated. Coordinates have been processed using different permanent stations (BME, Penc, Eszék, Graz, Bratislava), with code- or phase-method, with precise- or broadcast ephemerides, using measurements with 1-5-10-30 minute duration, observed on 4-5-6-8 satellites. We have measured the caves both in the summer, when the vegetation were high, and in the winter, when only the stocks and arms of the plants and the hang over rocks were covering the sky. We have investigated the loss of lock, loss of satellites near the horizon or the zenith.



Figure 1. Determination of a cave entrance

The short summary of the results:

1. Although GNSS satellites at the entrance of the caves are often highly masked, measurements with the Leica and Trimble GPS could also be performed. The accuracy of the navigation receivers were not sufficient, since those instruments can only provide reliable data if they have been initialized beforehand.
2. The optimal measuring time was 10 minutes. It was found that the best solution can be achieved by processing the coordinates from the nearest permanent station. Processing with the phase-method or with precise ephemerides is unnecessary; it is sufficiently adequate to calculate the coordinates with the code-method and with broadcast ephemerides.
3. It is useful to set out the receiver as high as possible, because of the masking effect of trees and hang over rocks. The investigations were demonstrating that masking the satellites near the

horizon results in less errors, than masking of some satellites near the zenith. The most important effect of masking is due to the stocks and arms of the plants, while the season-dependent leaves have negligible effect on the accuracy.

4. It is sufficient to use a general transformation parameter file for transforming between WGS-84 and the Hungarian EOV for the whole area of Hungary. Also, use of some troposphere- or ionosphere-model was found to be not necessary.

III. SURVEYING THE UNDERGROUND WORLD

Traditional surveying instruments (compass, tape) are the widely used instruments in cave surveying. The surveying-method is in most cases the traverse surveying and the polar method. After the measurement the data should be adjusted and the corresponding error estimated derived. Nowadays, total stations and laser scanners are the most up-to-date instruments to carry out precise measurements in large numbers. The laser scanners provide more thousand measurements per seconds, representing the wall of the cave with point cloud in an extremely fine resolution. The measured point cloud then can be imported in almost every desktop mapping software. Conversion of the three dimensional point cloud to a digital terrain model is only a question of the adequate knowledge of the software. In a subsequent step, database can be juxtaposed to the map, containing attribute data, for example the colour or health condition parameters of the dripstone, estimated lifetime. The accuracy of the laser scanners exceed the demanded accuracy of cave surveying; only their cost and size limits their spreading. A lot of scanners are enabled to be set up at a point, to measure the instrument height, to orient – these functions are basic conditions of the measuring in a local coordinate-system. When we have a correct transformation equations between the local and a national projection system –the later in Hungary is the EOV– all editing and post-processing steps can be performed in the national projection system (Fig.2).



Figure 2. A laser scanner in the Cseszneki-cave

In the last few years the surveying instruments are going through a heavy development. The total stations are always improving, their automatization is gradually getting into higher level, enabling their applicability for more colourful tasks on the spectrum of the surveying works. The development of the GNSS-receivers allows real-time measuring methods, so in the field we can get phase-processed coordinates in real time with centimetre accuracy. In the first times it was feasible by using an own base station in a reference point, but due to the development of the GNSS-infrastructure, this can be solved by using stations of the permanent station network (also called active GNSS-network). Recently, in the last 5-10 years Smart-Stations have been developed, with are combined total station and GNSS-receiver instruments, allowing convenient solution of complex surveying tasks. Cave surveying is basically a similar task to a survey of the interior of a building, but in the former case survey of the entrance of the cave and the nearby area of it is also included. Several relevant information on the evolution of the cave can be drawn geologist and geographers based on karst-shapes of the surface of the mountain. Determination of the direction of water flows, or the relative position of underground objects and surface karst-objects may give us information on the location of undiscovered cave-arms. Nowadays, the speleologists are using traditional surveying instruments. The directions are being measured with a compass, the distances with tapes. With this instrumentation and the applied methods the errors are usually above expectations, especially in case of big horizontal and vertical caves. The measurement of the nearby area of cave entrances has been failed due to financial reasons and time limitations. Compass and tape are not sufficiently precise for a reliable surveying. Several underground objects are undiscovered due to the lack of known connection between the extant cave-arms and the surface [1]. This is important, not only from the aspect of the natural conversation, but also from the aspect of tourism or healthcare, for example the Abaliget-cave in the Mecsek-mountain, or a Kórház-cave in a Bakony-mountain are very important for cure pulmonary diseases. The appropriate surveying instrument to deliver sufficiently accurate geometry of the cave and the mountain (for detecting cave-arms) is the Smart-Station (Fig.3).



Figure 3. The Smart-Station

With a Smart-Station we do not need control points at the cave-entrance, since by using the RTK-method we

can manage the control-point densification. Determination of the coordinates of a single point takes only a few minutes, so we can determine a large number of points, depending on the shape of the surface and the cutting-off angle. Based on these points, we can measure a regular network of detail points (in essence a set of three-dimensional points), which can be the input data of the post-processing and visualisation. After the control-point densification the GNSS receiver and the total station can be used separately, the former by mounting it to an antenna-pole. Thus, the measurement work can be done parallel at two different part of the area, which makes the measurement more efficient both from financial and temporal aspects. The aim of the surveying either can be a line-levelling, which yields a vertical model of the surface, or a Digital Terrain Model (DTM), which combines horizontal and vertical information in one model. When the cave is not too narrow, we can use the total station under the ground. The polygon-points and the cross-sections can be determined simultaneously in one workflow, since with the laser-distance meter we do not need a reflector tape or a reflector prism. Smart-Stations can also be programmed, which can control the measurements, and they can be performed in an automatized way like a robot-theodolite. We can get in shorter time much more information than before, making the measurements much more efficient in every aspects. Caves are classically displayed in cross and longitudinal sections, in horizontal map in a stretched sections. However, this new surveying method allows the illustration of the geometric information on an isometric or plastic map. Both maps are represented in a three-dimensional coordinate system, resulting in an impressive visualisation of the caves.

IV. MAP PRESENTATION

The most relevant documentation of a cave surveying campaign is the resulted map. It visualizes the spatial range of the cave in context to the surface. The map is an important tool for effective and purposeful explorations. A cave-map, like a geodetic map is showing the geometry in scale, ground plane after a generalization. As caves are three-dimensional objects, the visualization in two-dimensions is not that efficient; we favour the three-dimensional mapping methods, or plan-maps with many attribute data. The most cave-maps are sections, where the duct width on the map refers to the real width of the duct. When the cave is a vertical one, more sensible to display it in two or more layers projected onto a vertical plan. We can illustrate the cave with an isometric, axonometric method, where we replace the cave-arms with prisms, or by the hydrotermical caves with ball-chambers. When we want make a useful map, we should be informed on the maps, which are regularly used by speleologists (Fig.4).

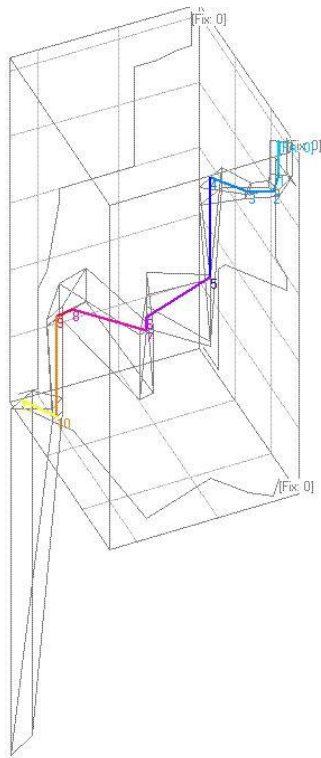


Figure 4. The 3D-model of the Wind-cave

V. BUILDING AN INFORMATION SYSTEM

The most effective, multi-faceted form of presenting the collected, spatial data of caves is the information system. It contains a lot of geometrical and attribute data, which is although a diversified database, can be practical by properly defined (simple or complex) queries (Fig. 5). Such information systems can be made available via the internet, so they can appease wide user claims ranging from the scientific to the touristic demands.

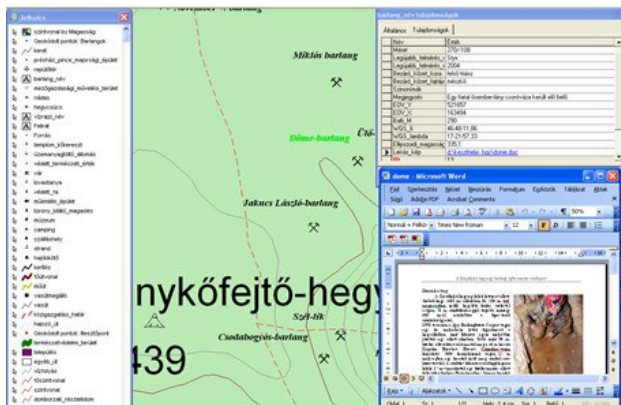


Figure 5. The information table of the Döme-cave

Building an Information System of caves, I can define two steps [4]. The first step is shaping the map, and the second step is compiling the descriptive attribute data. For shaping a map, it is very important to choose a good base map, with adequate content. Visualization of all important detail is demanded, when the representation of the actual geometry of the cave and of the surface is based on homogeneous point distribution. The most

suitable example is the Hungarian topographical maps in a scale of 1:25 000, which can conveniently be handled, and the details are presented in reasonable resolution. After choosing the adequate map-section, we must scan it with a professional scanner. For that either a flat-bed plotter or a drum-plotter can be used, though the latter is preferred by the author. In the drum-plotter not the sensor moves over the paper, but the paper is moved by numerous rolls. There is no positional error, and the plotter requires only a few places. Finishing the map-design, the next important step is collecting the attribute data. The attribute data is the actual purpose of the Information System, as this is the useful content which is juxtaposed to geographical locations. An important point of view in my project was that the data collection was widened to numerous kinds of data in cities, towns, villages and natural sights; in a given region there are not only caves, but cultural and ethnographical values as well. The operative part of the documentation was filling the attribute tables of the caves. All table contains the following information: cadastre number of the cave, the name of the cave, length, depth, name of the researcher spelunker club, the year of the last research, type of the rock, age of the rock, other ethnographical name of the cave, comments, coordinates in the Hungarian system (EOV_Y, EOV_X), height above mean sea (Baltic-sea), WGS coordinates (WGS_fi, WGS_lambda, ellipsoidal height), description, pictures, maps. These information are in a hyper allusion, under the name of the caves. Out of the cave's data, in this Information System further useful attribute content is provided, such the descriptions of the cities, towns, villages and natural signs. With this data, the information system can be used even for planning of a family excursion.

REFERENCES

- [1] Kordos L.: „Magyarország barlangjai”, Gondolat Könyvkiadó, Budapest, p.326, 1984
- [2] Eszterhás I.: „Magyarország nemkarsztos baarlangjai”, Kézirat, Isztimér, p. 30, 2000
- [3] Tarsoly P.: „GPS alkalmazása barlangbejárataok helyének meghatározására”, szakdolgozat, Székesfehérvár, p. 94, 2002
- [4] Tarsoly P.: „Cave Information System”, XIII. FIG Congress, Munich, Germany, Oct. 8-13. 2006, Surveyors Reference Library, (http://www.fig.net/pub/fig2006/papers/ts47/ts47_01_tarsoly_0340.pdf - accessed on 21.08.2015)

XML-based development of electronic library for teaching and educational subsidiary materials

Tsvetanka Georgieva-Trifonova* and Gabriela Chotova**

* Department of Mathematics and Informatics, University of Veliko Tarnovo, Veliko Tarnovo, Bulgaria

** Professional School of Electrical and Electronics, Gorna Oryahovitsa, Bulgaria

cv.georgieva@uni-vt.bg, chotovagm@abv.bg

Abstract— Recently, there has been increasing interest in the use of electronic libraries for teaching and educational subsidiary materials on informatics and information technologies at Bulgarian schools. There are noticeable possibilities for development and improvement of the existing systems, which would help to maximize the benefits of their applying. One such possibility is collecting of the information concerning the count and ratings of downloaded materials and their analysis by subjects, authors, keywords, type, etc. Similar analysis could assist to improve the quality of the materials in a way that they to be useful, as well as to arouse an interest. In the present paper, XML-based solution of that option is proposed, based on the advantages provided by the semi-structured data model.

I. INTRODUCTION

Nowadays, the world of information technologies is faced with the challenge of constantly growing software systems. Business software, social networking, online trading systems, electronic libraries generate huge amounts of data.

If we look back in time, database management systems had begun their way from automation tasks to record structured information. They go through several stages of development to the emergence of the concept of the relational model of Edgar F. Codd in 1970 [1]. This had made data modeling and creating software for their management much easier. Flexible structure, theoretical justification, standardized language and well-defined rules for construction are just some of the advantages of the relational model.

Beyond all other benefits, the relational model is suitable for systems of client-server type and has established itself as a key technology for storing and processing data of the web-based and business applications. But over time, these systems are becoming larger and are accessible to more users. With the entrance of Web 2.0 and the dissemination of social networks, in the foreground, the requirements for high performance, scalability, reliability and flexibility are placed. The relational databases management systems can be easily scaled vertically by adding supplementary hardware resources. However, the horizontal scaling, the delay caused by the excessive normalization and the increasing amounts of data are becoming a problem.

To solve these problems, much effort is exerted concerning the development of different types of non-relational (NoSQL) database management systems. They are used in everyday writing of comments by millions of people on social networks (Facebook, Twitter, Instagram,

etc.), data products in major electronic stores (Amazon, eBay), Google search and systems of large corporate funds.

The main problems facing most NoSQL databases are associated with [2]:

- data portability;
- compatibility between NoSQL databases;
- dependence on the provider (vendor lock-in).

XML databases overcome these problems as they take advantage of the standard languages: XML [3], query language XQuery, language for modifying XML data – XQuery Update (W3C Recommendations).

It is no accident that in order to solve these problems, some NoSQL databases offer integrated solutions such as integration of Zorba XQuery Processor with MongoDB, which is documentary-oriented and is currently the most popular NoSQL databases management system.

The advantages of building digital libraries based on XML databases can be summarized as follows [4, 5, 6]:

- It makes easier sharing of various resources;
- XML data are essential in a need of linked data, which are becoming more popular;
- The scheme of the data can be readily changed;
- It is possible to retrieve the information from data sources that are not restricted by a predetermined scheme;
- It allows flexible format for exchange of data between different databases.

II. MOTIVATION FOR USING XML DATABASE WHEN BUILDING A DIGITAL LIBRARY FOR TEACHING AND EDUCATIONAL SUBSIDIARY MATERIALS

In [7], we have examined current status of the implemented web-based multimedia information systems for e-learning materials and their use by students at Bulgarian schools. Besides, we have considered the role of electronic library in teaching information technologies at Bulgarian schools in order to provide more time for the application of the studied material and to increase the effectiveness of the educational process. There have been summarized the advantages and disadvantages of the use of digital libraries in teaching information technologies and are presented the main features of developed electronic library for teaching and educational subsidiary materials. The electronic library is built by means of HTML, CSS, PHP, MySQL.

After the conducted examinations of the current state of the creation of web-based multimedia information systems for e-learning materials and their use by students, it has been concluded that there is an outlining of the enhancing of their use in the learning process. The analysis shows possibilities for development and improvements to existing systems, which would help to maximize the benefits of their application.

Data storage for the materials in the database provides opportunities to collect information concerning the number and estimates of downloaded materials and their analysis in subjects, authors, keywords, type, etc. Such an analysis could help to improve the quality of the materials in order they to be useful and to arouse interest. Besides writing of comments to the electronic materials provided, it would facilitate the sharing of additional information (comments, additions, interpretations, clarifications) between users and authors of the materials.

Providing a possibility to maintain data on the number of downloads, ratings and comments on materials, is the main motivation for using the benefits of non-relational databases management system in the implementation of an electronic library for teaching and educational subsidiary materials.

Non-relational database system management is used for storage of data about the electronic materials together with the comments to them. Because the developed system must support search descriptions of electronic materials, the optimum is use of XML-based system for managing non-relational databases, since the built-in form of queries XPath / XQuery is standardized and optimized.

One of the fastest and lightest for using non-relational database management systems is BaseX. It is applied in a large number of software products – from those related to scientific activities, to libraries of XML documents. One of the customers of BaseX is Nexoma, developing platforms for Internet commerce and online applications. The software product which uses BaseX, is XS-KatalogExpress and it represents an electronic catalog of products.

These facts determine the use of BaseX in this present system to maintain electronic materials as an appropriate XML database management system.

III. DEVELOPMENT OF XML-BASED ELECTRONIC LIBRARY FOR TEACHING AND EDUCATIONAL SUBSIDIARY MATERIALS

In the present study, we transform the existing relational model, described in [7], in a model of semi-structured data (Fig. 1a).

The successors of the root element `elibrary` are:

- `materials` with one or more successors `material` with attributes `materialid` and `authoridrefs`, referencing one or more elements `author`;
- `authors` with one or more successors to the `author` with attributes `authorid` and `materialidrefs`, referencing one or more elements `material`.

The successors of the element `material` serve to present the characteristics of a given material – title (`title`), subject (`subject`), category (`materialcategory`), description (`description`), keywords (`keywords`), files (`materialfiles`), as well as additional – estimates (ratings) and comments (`comments`). The element `author` has successors intended to describe authors – name (`firstname`), family (`lastname`), title (`title`), short biography (`shortbiography`), username (`username`) and password (`password`).

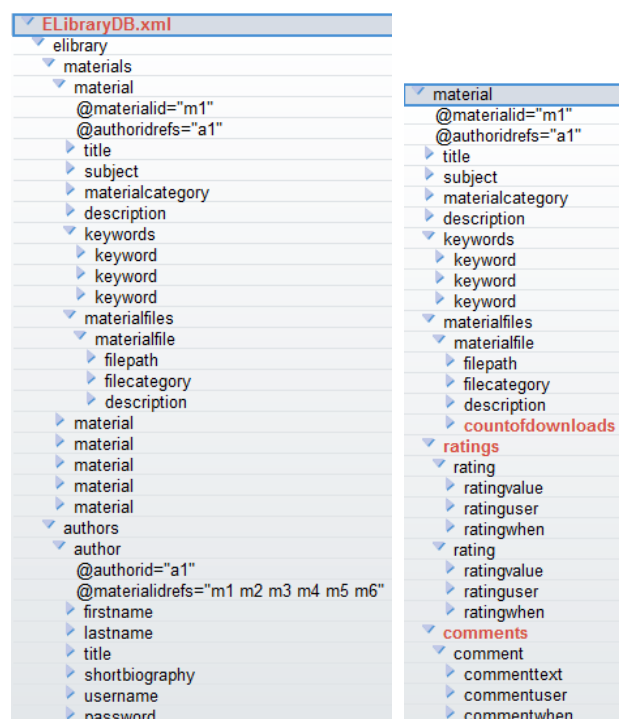


Figure 1. XML database ELibraryDB: a) Transforming the relational model into semi-structured data model; b) Extension with elements for downloading number of a file, comments, and material's ratings

Moreover, we have extended the resulting model with the following elements (Fig. 1b):

- Counting the file downloadings (`countofdownloads`);
- Writing the comments – only for registered users (`comments`);
- Evaluation of materials – only for registered users (`ratings`).

For implementation of the database ELibraryDB has been used the XML database management system BaseX (<http://basex.org>).

There has been set a type of document by DTD (Document Type Definition) and XML Scheme. Modeled and implemented XML database ElibraryDB is validated in accordance with:

- the definition, the definition, stated in <http://practicum.host22.com/data/elibrary/ELibraryDB.dtd>;
- XML scheme, the definition, stated in <http://practicum.host22.com/data/elibrary/ELibraryDB.xsd>.

Figure 2 shows the visualization of the XML scheme of the database ELibraryDB, using the online editor (XML Editor / Viewer Online) xmlGrid.net, accessible from

<http://xmlgrid.net>. Figure 3 shows the element material in more details.

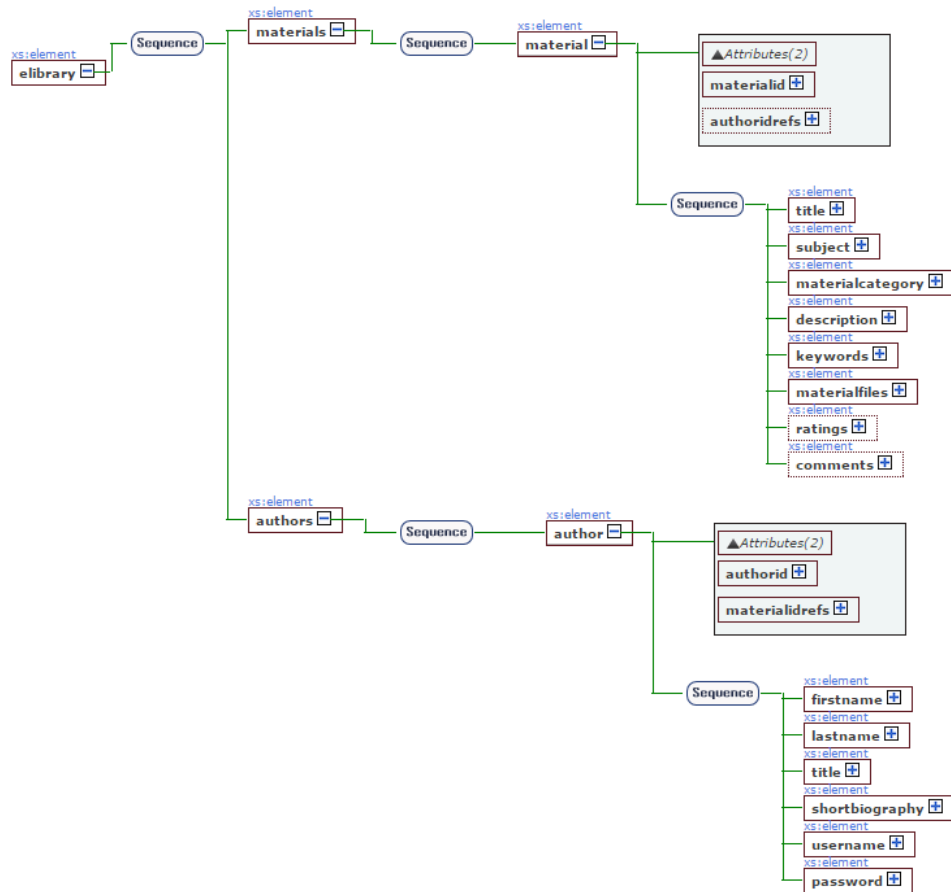


Figure 2. XML scheme of the database ELibraryDB, visualized with xmlGrid.net

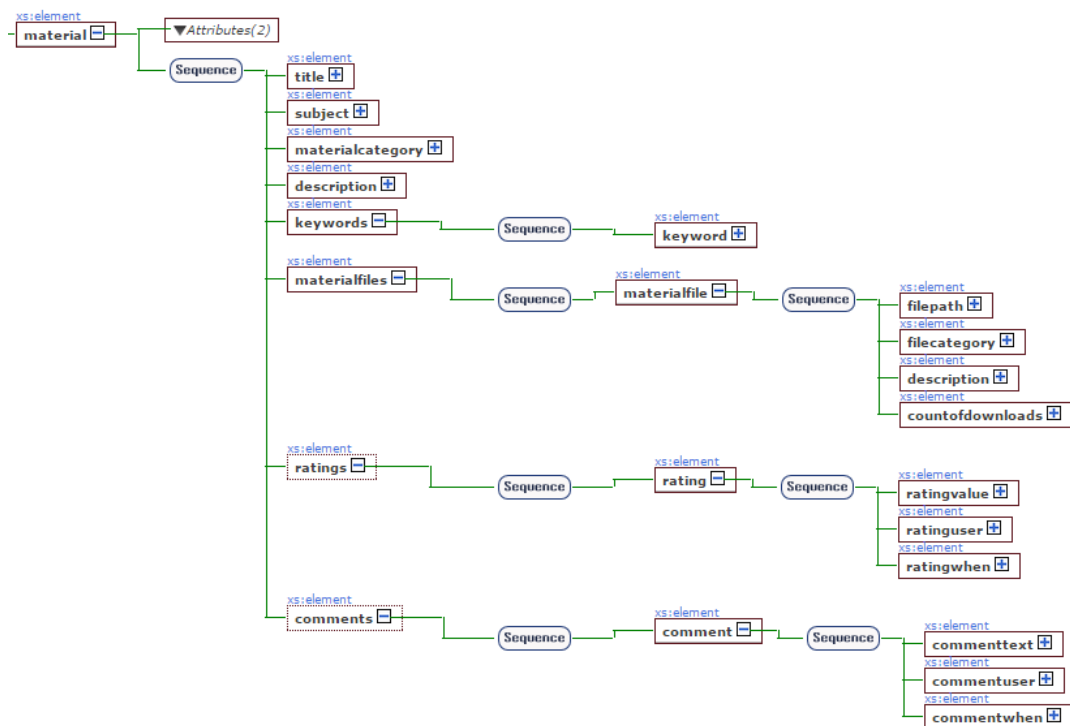


Figure 3. Successors of the element material

By using queries of XQuery, different requests needed for the given application can be retrieved. For example, a query of XQuery to retrieve data about materials on a given subject (Fig. 4):

- Title;
- Category;
- Description;
- Keywords;
- Files – location, category, description, number of downloads;
- The average of ratings by customers;
- Number of users who have given evaluations;
- Comments – text, user, when the comment is written;
- Authors – names.

```
<result>
{
  for $x in
    /elibrary/materials/material
  where $x/subject = "Database management system"
  return
    <material>
      {$x/title}
      {$x/materialcategory}
      {$x/description}
      {$x/keywords}
      {$x/materialfiles}
      <avgrating>
        {fn:avg($x/ratings/rating/ratingvalue)}
      </avgrating>
      <countrating>
        {fn:count($x/ratings/rating/ratingvalue)}
      </countrating>
      {$x/comments}
      <author>
        {fn:id($x/@authoridrefs)/firstname}
        {fn:id($x/@authoridrefs)/lastname}
      </author>
    </material>
}
</result>
```

Figure 4. Query of XQuery for retrieving data about materials on a given subject

Our future work includes creating of a convenient web-based user interface by means of HTML, CSS, PHP.

IV. CONCLUSION

In this paper, we have proposed XML-based building of an electronic library for teaching and educational subsidiary materials. The existing relational model is transformed into a model of semi-structured data and is extended so that it is possible to store the number of downloads of materials, writing comments, evaluation of materials. On the one hand, this additional information can be directly used by users, on the other hand it can be analyzed and the results of the carried analysis can be applied to improve the quality of materials.

One of our future tasks is to design and to develop an ontology of linked data by applying semantic web technologies.

REFERENCES

- [1] E. Codd, A Relational Model of Data for Large Shared Data Banks, *Communications of the ACM*, vol. 13, No. 6, 1970, pp. 377-387
- [2] D. McCreary and A. Kelly, *Making Sense of NoSQL: A guide for managers and the rest of us*, Manning Publications, 2013
- [3] K. Williams, M. Brundage, P. Dengler, J. Gabriel, A. Hoskinson, M. Kay, T. Maxwell, M. Ochoa, J. Papa, and M. Vanmane, *Professional XML Databases*, Wrox Press, 2000
- [4] K. Banerjee, How Does XML Help Libraries?, *Computers in Libraries*, vol. 22, No. 8, 2002, <http://www.infoday.com/cilmag/sep02/banerjee.htm>
- [5] L. Qiu and Y. Peng, Research on XML-based component library and its application, *Computational Intelligence and Industrial Applications*, 2009
- [6] R. Tennant, ed., *XML in Libraries*, New York, Neal-Schuman Publishers, 2002, 175 pp.
- [7] T. Georgieva-Trifonova and G. Chotova, Development and role of electronic library in information technology teaching in Bulgarian schools, *International Journal of Advanced Computer Science and Applications (IJACSA)*, vol. 6, No 3, 2015, pages 53-58

First results of implementing satellite-borne gravity data to GIS and future perspectives

L. Földváry^{*,**}, M.I. Kemény^{***} and X. Z. Huang^{****}

^{*} Institute of Geoinformatics, Alba Regia Faculty of Engineering, Óbuda University, Székesfehérvár, Hungary

^{**} MTA-CSFK Geodetic and Geophysical Institute, Sopron, Hungary

^{***} Department of Geodesy and Surveying, Budapest University of Technology and Economics, Budapest, Hungary

^{****} Institute of Remote Sensing and Digital Earth, Chinese Academy of Science, Beijing, China
foldvary.lorant@amk.uni-obuda.hu, kemeny.marton@epito.bme.hu, 380651945@qq.com

Abstract—Gravimetry has not acquired GIS for its own benefits. The reason is the basic difference in tools and demands. The global gravity field is typically described analytically by an infinite set of nearly ellipsoidal layers, which are synthesized to actual values afterwards. In contrary, GIS employs pre-defined data sheets, which inhibits the flexibility of the use of spherical harmonics. The present study discusses how gravity field information can be implemented in GIS keeping its basic approach to data sheets. As a result, a gravity module for a remote sensing GIS, termed Gravity_RS_GIS has been designed and developed.

I. INTRODUCTION

Geographic information systems (GIS) are basic tools for representation, manipulation and analysis of spatial or geographical data. In practice, the strength of GIS with respect to classical map-based representation of location-based quantities is that they are supported by interactive tools, i.e. queries can be defined by the users, and different spatial data analyses can be implemented, even the map content according to the result of the analyses can be edited. All in all, GIS provides a much more flexible platform than electronic maps, therefore they became popular for numerous applications ranging from the simplest visual data screening through location intelligence studies to severe, elaborated geoscientific research.

It is also frequently used for remote sensing data. For such a case a GIS provides a frame, where remote sensing observations from different sources can be integrated into one system, in which location-wise data management can efficiently be done. In contrary, satellite gravimetry has not really acquired GIS for its own benefits. The gravity field is typically a space-dependent physical variable, which is usually described in geosciences by analytical description of the shape of its equipotential surfaces, which is equivalent to an infinite set of nearly ellipsoidal layers [1]. Mathematically these analytical surfaces are described with spherical or ellipsoidal harmonics.

Even though, the GIS era has been started in the late 70s or 80s, its potential has not become clear for the gravimetry community even until the 90s. Among the first papers on use of GIS for the gravity field, some has been inappropriately addressed, e.g. [2] has labeled an analysis on the role of the marine gravity in measuring the sea surface topography as 'GIS data sets optimization'. An expedient attempt to test the applicability of GIS for geodesy has been delivered by [3]. This study has used the

GRASS GIS for integrating and processing a geoid model and satellite altimetry data. Basically, the GIS tools, which were actually applied in that study were 1) representation of the data sets (both gridded and sparse data), 2) interpolation of one data to the points of the other, 3) outlier detection by comparison the interpolated data with the other, 4) smoothing of the outlier-eliminated data and interpolation back to the original points. Even though they have concluded that GIS was fast and reliable, they have also found that it may have a little, but meaningful impact on geodesy. By the time, both GIS tools and geodetic observation techniques improved a lot, as so, it makes sense to revisit this conclusion.

An essential progress in geodesy in the last two decades is the globally available high-resolution gravity field models due to the dedicated gravity field missions, CHAMP, GRACE and GOCE [4]. Furthermore, due to these satellite missions temporal variations of the gravity field have become detectable globally on monthly scale. Large scale monitoring of gravity has opened a wide range of applications in geosciences, since every geophysical phenomena can make use of this information, which produce mass variations over a large area on long (i.e. longer than some month) time scales. Most of such mass variations occurs in the fluid envelop of the Earth (also referred as geosphere), which is a 'thin' layer on and above the surface of the Earth enclosing the masses of the atmosphere, hydrological processes, the world ocean and the cryosphere. Considering the effects of these developments, in the present paper the utility of GIS for satellite gravimetric applications is addressed.

II. BENEFITS OF GRAVITY INFORMATION FOR REMOTE SENSING GIS AND VICE VERSA

Processing of satellite gravimetric observations to deliver a gravity field model is a laborious task. It does not provide unique solution, as it is mathematically an inverse problem, so no routinely applicable processing method of gravity satellites observations can be defined. Instead of using GIS for gravity satellite data processing, its resulted gravity field models can more efficiently be used for further applications within a GIS software. Such applications can make benefit of different properties of gravity field; here are three examples.

A. Gravity field appoints horizontal and vertical directions

Question of height systems: most Earth-observed satellite-borne data are positioned in WGS-84 ellipsoidal

reference frame as (ϕ, λ, h) . The data can be transformed to local or regional coordinate systems, which is already implemented in every GIS software (either the latter being already defined in the software or by enabling to load its transformation coefficients in). For large regions, global scale analyses, calculations are performed in the ellipsoidal geographical coordinate system. However, ellipsoidal geographical coordinate systems are not connected to the horizontal and vertical directions, i.e. two points with exactly the same elevation above an ellipsoid can be at different heights. In true geographic coordinates same elevation values defined by their potential difference presented in Geopotential Unit (GPU) refer to the same height. The difference of ellipsoidal and true geographic heights is described by the geoid undulation, which can reach approximately ± 100 m. If geoid undulation data is available worldwide with appropriately fine resolution, than the ellipsoidal to true geographical coordinate system transformation, i.e. (ϕ, λ, h) to (Φ, A, W) can be defined. It can be substantial if an application is extended over a large region.

B. Gravity field manifests mass distribution

As the gravity is generated by mass, gravity anomaly provides information on the mass distribution of the underlying geological features as well. Gravity anomaly can be used as key tool for mineral exploration and geophysical prospecting, even though nowadays the more efficient seismic methods are preferred over time consuming gravimetry. A new role of gravity in prospecting may arise with the development of the satellite gravimetry, as they provide gravity information in a gradually increasing resolution. In the near future it may reach so fine resolution that the test area selection can efficiently be done based on satellite gravity data only. With this, the field work can be reduced notably.

C. Gravity field is a force field

The knowledge of the gravity field in fine resolution in space may be useful for aeronautics. Motion of a missile, an aircraft or any other flying object in the gravity field can be precisely predicted by taking into account the effect of the propulsion force and the gravity field. Such orbit integration need may arise for several segments of engineering ranging from telecommunication industry to military applications.

III. GIS APPLICATIONS FOR GRAVITY SO FAR

There are very few published attempts for applying GIS softwares for gravity field variables are known so far. An online GIS service, the Gravity Information System (SIS) has been developed at the Physikalisch-Technische Bundesanstalt (PTB), Braunschweig, Germany [5]. This system is presented on a Java platform, which enables query of gravity information (the gravity value, or free-air or Bouguer anomaly, also contour lines of gravity referred to as 'gravity zones') by coordinates (or by clicking on the base map) at the physical surface (approximated by the SRTM DEM model) or at any arbitrarily defined altitude above it. The background of the workspace can either be topographic or gravity anomaly map. This system makes only partially use of the benefits of the GIS in both on data query and on data representation.

The China Regional Gravity Information System (RGIS) has been developed by the China Geological

Survey [6]. This system is based on the MapInfo platform with OLE technology. As it is stated in [6], the system can visually interpret data of spatial geography, geology and gravity for China. It furthermore enables graphical data editing and data table operations, and what is really specific is that there are classical gravimetric tools are defined, such as gravity reduction, gravity (and magnetic) field transformation, and gravity anomaly inversion. The system has been developed for regional use. The system is developed for terrestrial gravimetric data. Unfortunately, the system cannot be tested, as no internet availability could be found.

Beyond these two essential examples, few other systems are known world-wide, [7], [8], [9], [10]. However, the abovementioned two examples are sufficient for demonstrating the status of the use of GIS for gravity data: There are some freely accessible softwares, which only partially make benefit of the GIS, and there are others, which are probably much more delicate (usually for regional use), however, these cannot be accessed for common users. In the present study a compromise, a trade-off is suggested: to deliver a GIS software for global use with features developed for both regular and geoscientist users.

IV. EXPECTATIONS WITH A GRAVITY GIS SOFTWARE

In fact, the convenient form of storing data for a GIS is providing them in tables, on data sheets. Therefore all gravity related quantities from a state-of-the-art global gravity model should be determined before hand, and include in gridded form to the GIS. It, however, makes the GIS software inflexible, which may have unlike effects in the long run. It should be remarked that global gravity models are developing quickly in the recent decades, c.f. the tables of global gravity field models at <http://icgem.gfz-potsdam.de/ICGEM/modelstab.html> [11]. Therefore, to stick to one gravity model is not wise. Furthermore, it excludes the use of local or regional gravity models. Users of the GIS software may have the option to load in local gravity field models according to their requirements and define its eligibility area. This would satisfy a more realistic need, since small features cannot be described by a global model, but a locally optimized one, as it is actually done in the practice of the different state land surveying entities. As local models may be undisclosed to the public, the need of integrating own model into the software is realistic. Therefore, we suggest supplying an appropriate tool for conversion gravity field models to gridded variables. In the present section, feasibility of inclusion of such a module inside the GIS is analyzed.

Basically, the mathematical tools for representing a gravity field model depend on the size of the area of interest. For global scale, the gravity is represented by analytical surfaces using the spherical or ellipsoidal harmonics. For local or regional scales, a gravity model can be equivalent to a set of point values, which are assumed to be sufficiently dense for reliable description between the points by linear interpolation. Analytical functions for regional models are also often defined in practice. In fact, spherical harmonic representation is so convenient and wide-spread that often they are applied for local gravity field models as well, even though those coefficients are valid only within the area of interest, and provides nonsense gravity properties outside of it. Local

or regional gravity field models can also be presented in a form of locally defined base functions, such as the spherical radial base functions [12], or by appropriately chosen high order polynomials.

All in all, the spherical harmonic representation is a feasible and regularly applied method, probably the most common tool. As so, a GIS containing gravity information is suggested to be prepared for including gravity field model in spherical harmonic form.

As for a GIS software, zooming in and out of the map is a basic tool, however, the complexity of describing the gravity field is notably different by the scale. Due to zooming into the map, a better spatial resolution should be defined automatically, which increases the number of computational points, e.g. by halving the grid size, the computational load is four times larger. This can be cured by fixing the size of the window, i.e. fixing the number of computational points regardless the extent of zooming. Disregarding this option, in Figure 1 the computational load by increasing the refinement of the grid is displayed with blue color.

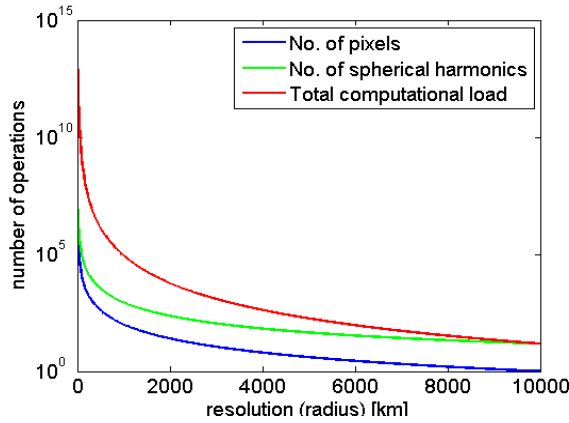


Figure 1. An estimate of computational load by determining the number of operations as function of the resolution of a grid.

The more essential consequence is that gravity information content should also be refined according to the extent of zooming. When using spherical harmonic synthesis for determining gravity from model coefficients, the spatial resolution is controlled by the choice of the maximal degree of spherical harmonics. The resolution is inversely proportional to the maximal degree of spherical harmonics, approximately the radius of a detectable gravity feature on the surface can be determined as

$$r = \frac{2R\pi}{l_{max}} \quad (1),$$

where R is the mean radius of the Earth and l_{max} is the maximal degree of spherical harmonics. According to (1) the computational load by increasing the maximal degree by 1 is increased by

$$\frac{(l_{max}+1)(l_{max}+2)}{2} \quad (2).$$

Fig 1 shows that how the computational load is increasing by increasing the maximal degree of the spherical harmonic expansion with the green curve. The total increase of the computational load, i.e. the sum of the

two is displayed with the red curve, showing that by drastically refining the resolution, the computational load may turn to a massive burden; no real-time calculation of gravity is suggested. Or, another option is that real time calculations are only performed occasionally for limited area.

Summarily, the GIS should provide gravity field information in a computationally effective way: the data should be 'quickly' accessible, still precise and up-to-date. Because of this, the inclusion of spherical harmonic synthesis is disregarded. According to that and the previous sections, two scenarios for gravity-involved-GIS are suggested:

1) "Data sheet scenario": Store all gravity related quantities with high resolution calculated up to the available maximal degree of spherical harmonics in data sheets from a global gravity field model, and the GIS software would handle only those data. When the user zooms into too small regions, in order to avoid pixelisation, the fine resolution can be achieved by a 2D spline interpolation.

2) "Semi-active scenario": Store all gravity related quantities with high resolution calculated up to the available maximal degree of spherical harmonics in data sheets from a global gravity field model, and the GIS software would handle only those data. This would serve for the small scale usage, i.e. for large area. When the user zooms into smaller regions, after a limit the software may aware that preferable local gravity model is available over the global one, or if not available, offer an upload of a local model. With an uploaded local model the spherical harmonic synthesis would be performed immediately.

With the use of the latter scenario, the users are more aware of possible choices of gravity models, therefore all parameterization would be done more cautiously. By offering a real time upload of gravity models, and performing the spherical harmonic synthesis real time, the computational load may notably increase. However, as this is suggested for 'small' area only, with an appropriate choice of the grid size (which can be relatively rough), but including all available coefficients, the computational load can be decreased compared to the red curve of Fig. 1. Note that for user uploaded models the area of eligibility should also be defined, for which cannot be controlled automatically, the user is expected the mark the borders.

Further alerts are suggested for both scenarios. These alerts are evident for users with satellite gravimetric background, but may cause a misapprehension by non-geodetic users. These are:

1) In order to keep up to date with the gravity field model, a spherical harmonic synthesis module can be included, which can be operated separately from the GIS functions, when no task is loaded in. Anytime, when the program is launched, the program should display the actual gravity model, and also should show a link to the most up-to-date gravity models for keeping the user informed on the most recent gravity models.

2) If temporal variations at a certain location are queried, a warning should inform the user whether new epochs of the gravity model time series are available. Also the link to the used gravity models' download page should be provided.

What gravity related quantities should be involved? As gravity data is considered to be integrated in GIS for non-

geodetic users of geosciences, the most useful quantities are probably geoid undulation and gravity anomaly, and for investigating temporal mass variations, surface mass anomaly. Optionally, gravity gradients, deflection of the vertical and geopotential can also be included in case of demand. All variables should be accessible on the physical surface and on the geoid as well. In order to provide these data at any height (outside the mass of the Earth), first and second vertical gradients of these quantities should also be stored in a data sheet, and then the approximate value of a variable at an elevation of h can be estimated by the first two terms of Taylor series of that variable.

If better accuracy is demanded, the use of spherical harmonic synthesis is unavoidable.

The core of the gravity field information in the GIS is the stored high resolution geoid undulation and gravity anomaly. In case of a query, the quantity of interest should be defined before hand, which may be switched to the other with a simple combination of keyboard buttons. When a quire is posed on them for a location, the GIS may pop up a window showing the vicinity of the point displaying the local structure of the quantity of interest, and write the actual value in that point. If the quire refers to an area defined by its corners (graphically selected or typed on keyboard), then the quantity of interest should be displayed for that area.

Temporal variations of the gravity field can also be queried for the selected location. For point-wise query, the GIS should display a time series of (GRACE-borne) surface mass anomalies. On the figure also a best-fit function (in a previously set form) may be plotted, which can be either a polynomial regression, a high-pass or low-pass filtering, or a periodic function.

Further, more elaborated, higher level tools for temporal gravity field variations can also be demanded. A linear trend or an annual variation may have more practical sense, if there are known unlike effects, and they are taken into account. For example, for ice mass balance studies linear trends are also affected by the motions of the tectonic plates, e.g. [13]. These motions are partially effects of the isostatic rebounding process since the relieve of the ice mass of the last glacial cycle, and also of the change of the ocean basins in this period [14]. For ice mass balance investigations, the GIS may offer a GIA models for correction of the linear trend, such as IJ05 [15], ICE6G [16], or W12a [17].

Surface mass anomalies may also contribute to determination of focal mechanism of large earthquakes afterwards [18]. By setting the date and the location of an earthquake, the extent of the change in the bias of the surface mass anomaly time series is informative. At any locations in the vicinity of an earthquake the time series of surface mass anomaly may be supplemented with bias and linear trend estimates before and after the event.

For application temporal gravity data for determining water mass being involved in the hydrological cycle of a basin, surface mass anomaly estimates should be limited to the area of the basin, and exclude masses from outside of the area to leak into the solution. For that, supplementary the shape of the basin should be outlined, and with the use of a hydrological model leakage correction should be applied as described by [19]. As this correction is a complex calculation, which is performed in

the spherical harmonic synthesis step, it is suggested to be performed real time in the "semi-active scenario".

The list of suggestions is incomplete, there may be several other tools invented, e.g. for oceanographic users. These should be delivered in case of demand.

V. THE GRAVITY_RS_GIS MODULE

To illustrate the effectiveness of the proposed framework for gravity application GIS software, a gravity module of the Remote Sensing GIS, termed Gravity_RS_GIS, is designed and developed based on the DotSpatial library. DotSpatial is an open-source geographic information system library, which allows developers to incorporate spatial data, analysis and mapping functionality into their applications [20]. Gravity_RS_GIS is characterized by the integration of actual gravity field data for further GIS applications or gravity field research. The present version of Gravity_RS_GIS is developed according to the "data sheet scenario".

The actual gravity field data basically includes geoid undulation, gravity anomaly and surface mass anomaly. For the gravity anomaly and the geoid undulation data the 2.5 x 2.5-minute resolution free-air gravity anomaly and geoid undulation grids of EGM2008 were used. [21]. As for the surface mass anomaly, GFZ RL05a GSM models [22] were used in the period of April 2002 to April 2015 up to degree and order 90. The models have been smoothed with a Gaussian filter of 500 km and the de-striping filter of [23] was applied. For more details on the methodology see [24].

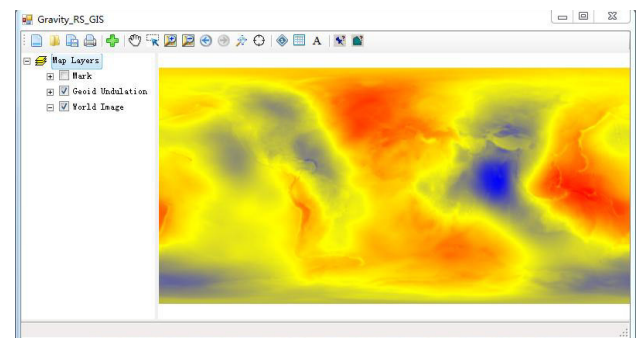


Figure 2. The geoid undulation data visualized in Gravity_RS_GIS.

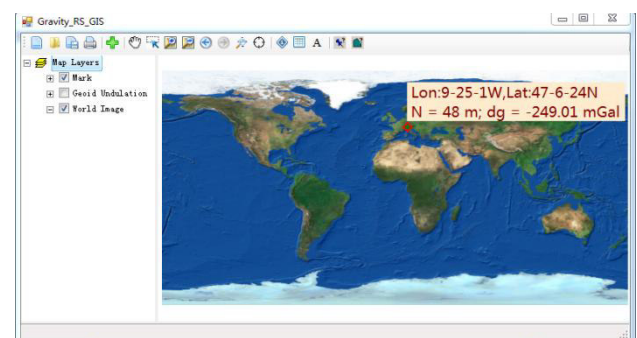


Figure 3. Gravity field data queried as back information after some remote sensing image in Gravity_RS_GIS.

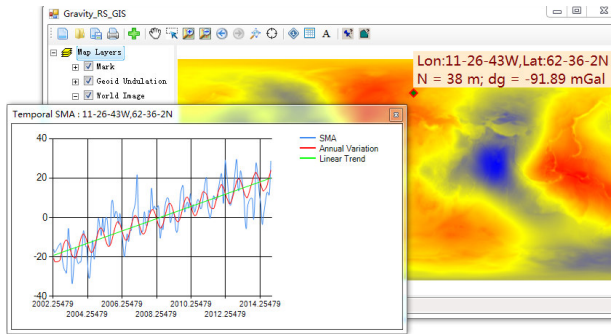


Figure 4. A time series query for temporal SMA data in Gravity_RS_GIS.

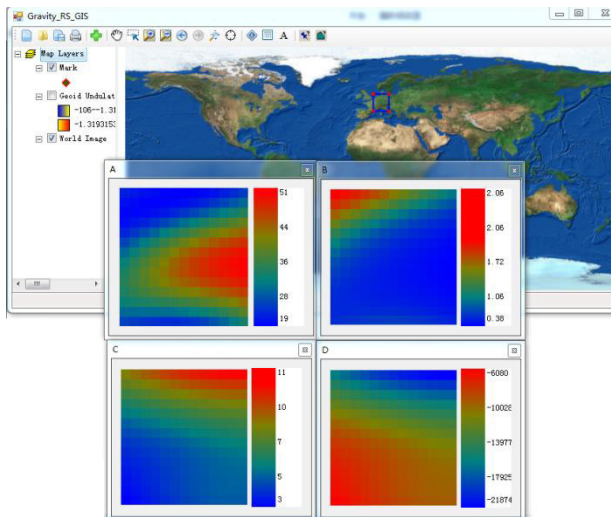


Figure 5. Temporal variations of the gravity field for a specific area in Gravity_RS_GIS.

The geoid undulation data is in a GIS raster data format, the others are not, therefore the data was reorganized and geo-referenced after loaded into Gravity_RS_GIS. And then, the data could be visualized, queried and analysed as gravity field information, as shown in the following figures. The basic features are as follows: (1) The geoid undulation data can be visualized in Gravity_RS_GIS (Fig 2) or as back information after some remote sensing image (Fig 3). (2) A time series query for temporal surface mass anomaly data in Gravity_RS_GIS (Fig 4). (3) Temporal variations of surface mass anomaly for a specific area described by the annual amplitude, annual phase shift, the linear trend and a bias (Fig 5).

VI. FUTURE PERSPECTIVES

For later versions, the implementation of the “semi-active scenario” is initiated. For the purpose, the cornerstone is a spherical harmonic solver, which can computationally efficiently, real time perform the spherical harmonic synthesis step. Parallel to the Gravity_RS_GIS, such a software is also developed in C++ language under the name BME Spherical Harmonic Solver (BME-SHS) [25]. At the present stage, the BME-SHS can handle geopotential, gravity anomaly, gravity disturbance, geoid undulation, and the six independent gravity gradients. In order to integrate it into a later

version of Gravity_RS_GIS, surface mass variations should also be implemented. The BME-SHS delivers solution both for spherical harmonic synthesis and for spherical harmonic analysis. For the analysis, there are limitations in the applicable maximal degree of the spherical harmonics. This is, however, no concern for the Gravity_RS_GIS, as it make use of the synthesis only.

Even though at the moment, in the second half of 2015 the gravity database included in the Gravity_RS_GIS are definitely up to date, it should be noted that within some years certainly new versions of the GRACE gravity models will be presented, as the present models, produced by different data centers differ notably to each other [24]. Therefore, regular update of the data set is essential, in later releases as well.

VII. SUMMARY

According to the present practice and demands, GIS may not be a central tool for processing satellite geodetic observations; instead, results of satellite gravimetry may efficiently support other applications, such as RS-GIS. Gravity field may work therefore as a tool for GIS.

The core of the study is a theoretical discussion, a stream of thoughts on implementing gravity information to GIS. It is found that though GIS basically works with existing spatial data, for gravity field applications several flexible parameters are demanded, which would make a GIS software inefficient. Therefore, for implementation of the gravity field variables for GIS a “data sheet scenario” and a “semi-active scenario” is proposed. The first case is a classical GIS data base, while the latter is a GIS database supported with physical geodetic applications.

A first version of a Gravity_RS_GIS has been implemented, which has been developed upon the “data sheet scenario”. Therefore, it has not made use of all theoretical findings, though the applicable tools have been implemented successfully. In later releases, the implementation of the “semi-active scenario” is proposed.

ACKNOWLEDGMENT

This study has been supported by the Chinese Academy of Science (CAS) President’s International Fellowship Initiative (PIFI).

REFERENCES

- [1] Heiskanen, W.A. and H. Moritz, “Physical Geodesy”, W.H. Freeman and Co., San Francisco, 1967.
- [2] Maslov, I.A., The use of gravity for GIS data sets optimization, *ISPRS Archives*, Vol. XXXI, Part B3, p. 511-516, 1996.
- [3] Crippa, B., Sanso, F., A Quick Look on the Sea Surface Topography of the Mediterranean from the Geomed Geoid and the ERS 1 Geodetic Mission, *Proceedings of the international centre for mechanical sciences workshop on Data acquisition and analysis for multimedia GIS* (eds. Mussio, L., Forlani, G., Crosilla, F.), Springer-Verlag New York, Inc. New York, ISBN:3-211-82806-0, p. 237-247, 1996.
- [4] Balmino, G., Perosanz, F., Rummel, R., Sneeuw, N., Suenkel, H., CHAMP, GRACE and GOCE: Mission Concepts and Simulations, *Bollettino di Geofisica Teorica e Applicata*, 40, 3-4, 309-320, 2001.
- [5] “Schwere-Informationssystem”, online GIS service, available at <https://www.ptb.de/cms/ptb/fachabteilungen/abt1/fb-11/fb-11-sis.html> (last check in 2015.09.15.).

- [6] Zhang M.H., He H., Wang C.X., The launch of a large regional gravity information system in China, *Applied Geophysics* 8(2): 170-175., 2011.
- [7] Hobbs, S. F., Kimbell, S. F., Coats, J. S., and Fortey, N. J., BGS databases for mineral exploration: Status in 2000, *BGS Research Report RR/00/11*, British Geological Survey, 2000.
- [8] Hinze, W. J., Aiken, C., Brozena, J., Coakley, B., Dater, D., Guy Flanagan, G., Forsberg, R., Hildenbrand, T., Keller, R. G., Kellogg, J., Kucks, R., Lee, X., Mainville, A., Morin, R., Pilkington, M., Plouff, D., Ravat, D., Roman, D., Urrutia-Fucugauchi, J., Véronneau, M., Webring, M., and Daniel Winester, D., New standards for reducing gravity data: The North American gravity database, *Geophysics* 70, J25 – J32, 2005.
- [9] Wang, C. X., Zhang M.H., Co-development of MapInfo and Surfer in regional gravity information system, *Geophysical and Geochemical Exploration* 32(4): 445-447, 2008.
- [10] Tracey, R., and Nakamura, A., Complete Bouguer anomalies for the Australian National Gravity Database, *ASEG Extended Abstracts* 1: 1 – 3., 2010.
- [11] Barthelmes, F., Köhler, W., International Centre for Global Earth Models (ICGEM), *Journal of Geodesy*, 86(10), 932-934, 2012.
- [12] Tóth, Gy., Földváry, L., Updated Hungarian Gravity Field Solution Based on Fifth Generation GOCE Gravity Field Models, *Proceedings of the 5th International GOCE User Workshop* (ed. Ouwehand, L.), Noordwijk: ESA Communications - ESTEC, Paper p_g11, 2015.
- [13] Shum, C. K., Kuo, C., Guo, J., Role of Antarctic ice mass balances in present-day sea level change, *Polar Science*, doi:10.1016/j.polar.2008.05.004, 2008.
- [14] Peltier, W. R., Closure of the budget of global sea level rise over the GRACE era: the importance and magnitudes of the required corrections for global glacial isostatic adjustment, *Quaternary Science Reviews*, doi:10.1016/j.quascirev.2009.04.004, 2009.
- [15] Ivins, E. R., James, T. S., Wahr, J., Schrama, E. J. O., Landerer, F. W., Simon, K. M., Antarctic contribution to sea level rise observed by GRACE with improved GIA correction, *Journal of Geophysical Research: Solid Earth*, 118: 1-16, 2013.
- [16] Peltier, W.R., Argus, D.F. and Drummond, R., Space geodesy constrains ice-age terminal deglaciation: The global ICE-6G_C (VM5a) model. *J. Geophys. Res. Solid Earth*, 120: 450-487, 2015.
- [17] Whitehouse, P.L., Bentley, M.J., Milne, G.A., King, M.A. and Thomas ,I.D., A new glacial isostatic adjustment model for Antarctica: calibrated and tested using observations of relative sea-level change and present-day uplift rates. *Geophysical Journal International* 190: 1464-1482, 2012.
- [18] Chen, J.L., Wilson, C.R., Blankenship, D.D., Tapley, B.D., Antarctic Mass Change Rates From GRACE, *Geophys. Res. Lett.*, 33, L11502, 2006.
- [19] Longuevergne, L., Scanlon, B. R., Wilson, C. R. , GRACE hydrological estimates for small basins: Evaluating processing approaches on the High Plains Aquifer, USA, *Water Resources Research*, 46, 2010.
- [20] "DotSpatial", available at <http://dotspatial.codeplex.com/> (last check in 2015.09.15.).
- [21] Pavlis, N. K., Holmes, S. A., Kenyon, S. C. and Factor, J. K., The development and evaluation of the Earth Gravitational Model 2008 (EGM2008), *J. Geophys. Res.*, 117, B04406, 2012.
- [22] Dahle, Ch., Flechtner, F., Gruber, Ch., König, D., König, R., Michalak, G., Neumayer, K. H., GFZ GRACE Level-2 Processing Standards Document for Level-2 Product Release 0005, *Scientific Technical Report STR/12/02 – Data*, revised edition, GFZ Potsdam, January 2013.
- [23] Swenson, S., Wahr, J., Post-processing removal of correlated errors in GRACE data, *Geophysical Research Letters*, 33(8): L08402, 2006.
- [24] Földváry, L., Kiss, A., Uncertainty of GRACE-borne long periodic and secular ice mass variations in Antarctica, *Marine Geodesy*, accepted, 2015.
- [25] Kemény, M., Földváry, L., Számítástechnikai fejlesztés nagyméretű műholdas gravimetriai adatok feldolgozásra, *Geomatikai Közlemények*, submitted, 2015.

4 D Model of E-Learning

D. Valcheva, M.Todorova and O. Asenov

St. Cyril and St. Methodius University of Veliko Turnovo, Bulgaria

e-mail: donika_valcheva@abv.bg, marga_get@abv.bg, olegasenovpc@gmail.com

Abstract— the paper presents a 4 D model for increasing the personalization of the e-learning process. Its 4 directions (4Ds) are: Previous student competences, Individual student's needs and style, Progress and Meaningfulness. An approach for application of the model is also suggested. It uses preliminary processing and simulation of the teaching and learning process for proactive assessment and transformation of the existed e-learning content towards the individual student needs.

I. INTRODUCTION

Our modern life is strongly influenced by effects such as rapidly changing and developing information, technology-enhanced communication and information access, and new forms of production and services.

This situation requires individuals to adapt their skills and competencies. Consequently, educational objectives and societal expectations have changed significantly in recent years.

The effectiveness of e-learning depends on the quantity and quality of the teaching materials, the time needed to pass through a course and the results achieved at the end.

Serious problem in e-learning is the low level of personalization of the teaching and learning process. In the Internet space can be found countless courses in one and the same theme, presented in different way, with different level of usage of multimedia elements, directed to different user profiles, with different duration and complexity.

The user has the very difficult task – to find the most appropriate course for his/her learning style, basic knowledge and skills. This is not always possible, and even when the choice of an appropriate course is a fact, the chance the initial goal (gaining knowledge and skills in a given field) to be reached for a short time is really minimal.

The personalization in e-learning may be described as a set of procedures, approaches and techniques for providing to students learning content, which is most appropriate for their needs, interests, previous knowledge and skills and individual learning style. This problem is deeply discussed and analyzed in many publications [1,3,4,5,6, 9,10,11, 16,17].

Personalized Learning is the tailoring of pedagogy, curriculum and learning environments to meet the needs and learning styles of individual learners. Personalization is broader than just individualization or differentiation in that it affords the learner a degree of choice about what is learned, when it is learned and how it is learned [18,19].

The adaptation of educational content to the personal characteristics of the students (learner profile) is a

complex process and is a research field for many authors nowadays [12,13,14,15].

It requires:

- Collection of user data (interests, level of knowledge in a particular area, learning style, etc.). This can be done in the registration stage.
- Storage of the received data and user profile support.
- Accumulation of data for the learner behavior in interacting with the system and in collaborating with other learners (in which scientific area are his/her interest, how often he/she used a given course, what results did he/she achieve in the final assessment in a given course, what keywords he/she used most often in searching information in the system, in what thematic area he/she was involved in collaboration with other students, etc.).
- Collection and accumulation of group assessment in order to be obtained adequate information of how much the learning units are interested to the students, meet their needs and preferences and are useful for improving their level of competence and their successful application in future employment or other training.

II. 4D MODEL FOR E-LEARNING– DESCRIPTION AND MAIN CONCEPTS

The 4 Ds of the model are shown on fig.1:

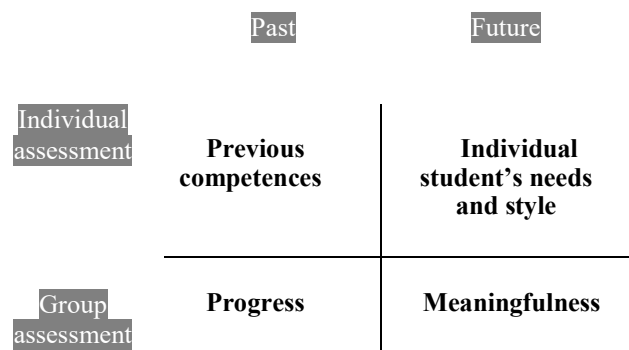


Figure 1. The concepts of the 4D model of e-learning

The model's 4Ds are: Previous student competences, Individual student's needs and style, Progress and Meaningfulness.

For the 4D model design are important the following main concepts:

A. Every student has individual needs and learning style.

There are a lot of theories for learning style and a lot of research could be found in this field [7,8,9,14,17]. In the present report is used the so called perceptual modalities or learning styles, defined according the sense perceptions.

There are three basic modalities to process information to memory: visual (learning by seeing), auditory (learning by hearing) and kinesthetic (learning by doing). Each of these different types of learning and absorbing the new information has specific characteristics.

The learning style reflects the way of absorbing and processing the information.

The human has the three styles but in different ratios. Every person develops preferred and consistent behavior and specific approaches to learning. This is related to three processes that form the differences in styles:

- knowledge – how the knowledge is acquired;
- conceptualization – how the information is processed;
- motivation and emotions – the way of making decisions and emotional preferences [2].

Many students do not know their dominant learning style. This is a very serious problem when searching for an appropriate for them learning material.

When the e-learning system “knows” the individual learning style of a given student, it could select and offer him/her the most appropriate learning content.

B. Every student has different previous competences, obtained in the past.

The level of competence of each student is different. It depends on many different factors – the knowledge obtained by previous passed courses, the individual needed time and speed of absorbing and processing the new material, the level of real interest in the subject area, the formalism and low level of personalization in the traditional form of learning – each curriculum consists of N disciplines distributed in K educational years, etc.

Regardless of the educational level and form of studying the curriculum and syllabus are preliminary defined. For obtaining the relevant qualification level the student has to pass successfully all disciplines, i.e. he has to gain all credits, according the ECST.

In the traditional form of studying, the teaching is realized in front of a group of students with different learning profile (style, needs, interests, individual characteristics, etc.).

In most of the cases, although the professionalism and the wish of the teachers to work individually with each student separately, the personalization is very low. The education period is fixed and is not depending on the speed of adoption of the new material of the different students.

C. Every course could be assessed by the students most precise at its end.

The assessment should be very useful if it is based on two different directions – the first one is the progress that the students detect after passing the course and the second one is the level of meaningfulness (their vision where and

how these competences could be applied for real problem solving or could help them understand more complex and specific knowledge). This could be easily realized by providing a survey among the students at the end of the course.

D. The group assessment implies greater reliability.

It is well known that sometimes the assessment of students may be not so realistic. It may be influenced by accidental negative for him/her event during the study, disparity in the characters of the course participants and as a result difficulties in the collaboration process, individual problems in understanding the material, etc. That is why the group assessment is so much important, because it ensures greater reliability and low level of subjectivity.

III. ONE APPROACH FOR APPLICATION OF THE 4D MODEL

The suggested in this section of the report approach applies the 4D model of e-learning. This approach ensures higher level of personalization of the e-learning by preliminary processing and simulation of the teaching and learning process for priori assessment of the effectiveness and for transformation of the existed e-learning content towards the individual student expectations.

The approach consists of seven main stages. They are visualized on fig. 2:

Stage 1 – the courses versions are description by experts according to the following criteria: for which learning style are most appropriate, level of complexity, what kind of previous competences are needed, for which thematic area are directed, group affiliation, etc. on this stage the course versions' profiles is created and saved in a Course profile database.

Stage 2 – includes the student's registration and determination of his individual learning style [1,7,8], previous competences and group affiliation. Thus the individual profile of the student is created and saved in a Student profile database.

Stage 3 – includes the student request for learning.

Stage 4 – software application, defines the best versions of the existed e-learning courses for the individual learning profile of each student. This choice is made in accordance with the specifics of each individual student profile (learning style, previous competences and group affiliation), the course versions' profiles and the results of the assessment of the group to which this student belongs (this includes his vision where and how these competences could be applied for real problem solving or could help him/her understand more complex and specific knowledge).

Stage 5 – the selected most effective e-learning course is offered to the student and the online activities are started.

Stage 6 – actualization of the student profile.

Stage 7 – this step includes student assessment about the progress that he/she detects after passing the course and the level of meaningfulness (his/her vision where and how these competences could be applied for real problem solving or could help him/her understand more complex and specific knowledge). The received information is saved in Group assessment database.

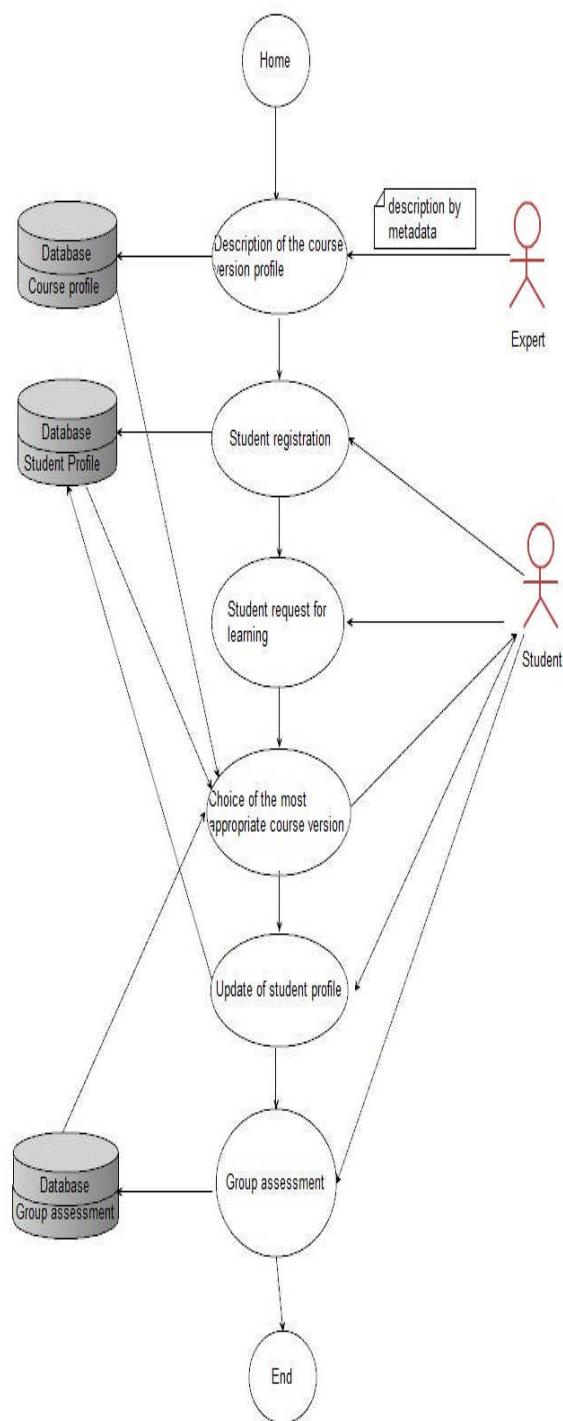


Figure 2. Approach for application of the 4D model

IV. CONCLUSIONS AND FUTURE WORK

In the future work the authors plan to define concrete criteria for determination of the group affiliations. The group affiliations are very important for the personalization of e-learning, because they define the needed competences within the group. It is also planned an experiment with current students from the two specialties “Computer Science” and “Informatics” of the VTU “St. Cyril and St. Methodius” and graduated students who are working already in companies to be provided and

analyzed in order the presented in this report 4D model and approach to be applied in real environment.

We plan to develop multimedia modules, which will be accessible by a cloud technology and each student will be able to learn according his/her own learning profile. For the pilot case we will choose 20 students, which will be grouped into 5 teams. The 4 members of each team will have one and a same group affiliation and will study and cooperate using the most appropriate for the group module. The results of this experiment will be analyzed and discussed. Thus a “group personalization” in the e-learning process will be achieved.

V. REFERENCES

- [1] Вълчева Д., М. Тодорова, Елементи на ефективното е-учене. Модул от платформа за е-обучение, позволяващ адаптиране на учебното съдържание според стила на учене, Списание “Computer Engineering”, София, бр.1/2010
- [2] Иванов И. Стили на учене. Втора национална научнопрактическа конференция “Психолого-педагогическа характеристика на детството”, Попово’ 2003, Университетско издателство “Св. Кл. Охридски”, 29-39.
- [3] Sants Arnaldo, Gomes P. Qualitative Evaluation of Student Participation in Distributed Learning Communities. 3rd E-learning conference, Coimbra, Portugal, 7-8 September, 2006, 2.17-1 – 2.17-6.
- [4] Schrum, L., Berge, Z. Creating student interaction within the educational experience: A challenge for online educators. Canadian Journal of Educational Communication, 1997, 26(3), 133-144
- [5] Sonwalkar, N. (2001). A New Methodology for Evaluation: The Pedagogical Rating of Online Courses, Syllabus. <http://www.syllabus.com/article.asp?id=5914>
- [6] Тотков Г., Р. Донев., Е. Сомова., Е. Шойкова., А. Ескенази., М. Райкова., В. Сивакова., Ст. Хаджиколева., Е. Хаджиколев., Д. Благоев., Св. Енков., Хр. Инджов., Д. Тупарова., Г. Тупаров., Р. Радев., А. Смрикаров., Д. Левтерова. Е-обучението в информационното общество. Университетско издателство „Паисий Хилендарски”. 2001. ISBN 978-954-423-651-9
- [7] Todorova M., T. Kalushkov, D. Valcheva (2008), INFLUENCE OF ICT ON LEARNING MODALITIES, EDU WORLD CONFERENCE, Pitesti, Romania
- [8] Valcheva, D., M. Todorova, T. Kalushkov, Structuring Multimedia Scenarios According To The Different Learning Modalities, EATIS 2009, Prague, Czech Republic
- [9] Wen Jia Rong; Yang Szu Min. The effects of learning style and flow experience on the effectiveness of e-learning. Advanced Learning Technologies, 2005. ICALT 2005. Fifth IEEE International Conference page. 802- 805.
- [10] Wickersham L., Barbara Tyler. Assessing the Quality and Effectiveness of Online Courses: A Qualitative Approach. <http://faculty.tamu-commerce.edu>.
- [11] Zhuhadar L. Romero E. Wyatt R. The Effectiveness of Personalization in Delivering E-learning Classes. Advances in Computer-Human Interactions, 2009, ACHI '09, Second International Conferences page 130-135, 1-7 Feb. 2009.
- [12] Sims, R. Interactive learning as an "emerging" technology: A reassessment of interactive and instructional design strategies. Australian Journal of Educational Technology, 1997, 13(1), 68-84.
- [13] Snajder Maja, Mateja Verlic, Petra Povalej, Matjaz Debevc. Pedagogical evaluation of e-learning courses – Adapted pedagogical index. Conference ICL2007 September 26-28, 2007 Villach, Austria.
- [14] Rong Wen Jia, Yang Szu Min. The Effects of Learning Style and Flow Experience on the Effectiveness of ELearning. Proceedings of the Fifth IEEE International Conference on Advanced Learning Technologies (ICALT'05) 0-7695-2338-2/05 \$20.00 © 2005 IEEE.
- [15] Rosenberg Marc. E-learning: Strategies for delivering Knowledge in the Digital Age. www.elearningeuropa.info/index.php?lng=1.

- [16] Kartir Lee. E-Learning: The Quest for Effectiveness. Malaysian Online Journal of Instructional Technology Vol. 2, No.2, pp 61-71 August 2005 ISSN: 1823-1144
- [17] Loo R. Kolb's learning styles and learning preferences: Is there a linkage?. Educational Psychology, 2004, 24 (4), 99 – 108.
- [18] <http://www.talentlms.com/elearning/personalization-and-learning>
- [19] Paneva D., Some Approaches for Personalization in Learning Management Systems, In D. Dochev, R. Pavlov (Eds.) "e-Learning solutions – On the Way to Ubiquitous Applications", Proceedings of the Joint KNOSOS-CHIRON Open Workshop, 26-27 May 2005, Sandanski, Bulgaria, pp. 65-74.

The role of the information system “Bulgarian army logistics” for material storage in Bulgarian armed forces

V. Banabakova *, S. Stoyanov *

*National Military University “Vasil Levski”, Department of Management and Logistics, Veliko Tarnovo, Bulgaria

e-mail: v_banabakova@abv.bg

Abstract— The information system “Logistics of the BA” (IS ‘Logistics’ of BA) is built in accordance with the Conception for Information Strategy of MOD referring to C4I systems ensuring the necessary compatibility with the systems of the other countries members of NATO. IS Logistics of BA is an integrated system. It is a complex and an instrument to manage and execute all the ordered movement with the material resources, including automated registering of all the accounts sheets in a previously set accounts module. The concrete benefits for the warehousing in BA are specifically in the spheres of: introduction of the Bar Code technology, and also in relation to the amount of data reports pertaining to the basic and additional parameters of the material resources- the object of safe-keeping.

I. INTRODUCTION

The information logistics presents an integral part of any logistical system, providing the functional area of the logistical management. The object of the information logistics are the information flows, reflecting the movement of the material, financial, and many others (operations) activities, related to the logistical support.

The information flow is generated by the material flow. Processing logistics information is the basis of material flow management. A necessary condition for the coordination of all the units of the logistical system, turns out to be the existence of an information system, which should be capable, to give, quickly and economically, the right signal, directed to the right point, at the right moment.

The existing logistics system of the Bulgarian Armed Forces includes management structures, logistical structures, infrastructure and transport communications, as well as, the the sequential order, forms and principles in all activities concerning logistical support, and finally, it also involves the communication and information systems.

In MOD and BA, they aim at developing a common informational environment, which has to be viewed as an integral part of the process of transformation. A part that will ensure the effective functioning of the command and management systems. [1] BA is already working on projects, concerning the automation of management activities and also aiming at the improvement of the informational support. Upon finalization of these projects, we will have a common information environment, and will be able to minimize time and resources spent in processing documentation. So far, what BA has developed and has been using, is based on the available segments of the public network ‘The Internet’ and the departmental

network “Automated Information System of the Bulgarian Army”(AIS of BA), which is used for sharing of classified information, giving opportunity for a number of common features (electronic mail, WEB, and other) to be used, as well as, giving an access to intranet departmental informational systems.

One of the main principles of logistics, is the requirement for clarity, as it is manifested through the access given, and by creating the automated system that gives account of the expenditures made, processes data and gives an easy way for sharing of information between the civilian and military structures in national and international environment of the logistics area of operation. [2]

Since 2011 in AIS of BA, we have the readiness to support an infrastructure using a public key, and maintaining all standard functions, based on digital certificates and the electronic signature.

Also a further development of the military messages exchange system by NATO standards has been planned, based on the standard electronic mail service in its full functional scope.

AIS will deal with estimate, estimate-analytical and reference information systems of individual and group information services of a large number of units of the BA. It encompasses more than 60 spots in the respective units. [3] The work on AIS is now focusing on the maintenance and exploitation of the already created infrastructure and subsystems on the one hand, and on the other - on its improvement and modernization. The existing infrastructure systems are partly sustained through outsourcing, on the basis of signed contracts with external firms.

The exploitation of the existing systems- the automated system for management ”Human resources”, the information system “Logistics”, “ Payroll department” “Documentation turnover” “ Information rubrics”, “ Operative archive”, “ Observation and control system’ and others, to a great extent, facilitate the administration of the processes, the management of information flows and the development and dissemination of accurate and timely analyses, that will assist in the taking of management decisions.

II. CHARACTERISTICS AND IMPORTANCE OF IS “LOGISTICS OF BA”.

The information system” Logistics” is an integral part and a subsystem to the logistical support in the armed forces.

One of the main requirements for the logistic system to be effective, is the existence of an information system, that will allow all the activities involved in the planning, contracting, supply, transport, storage and the rest of the logistical activities, to be combined together and managed according to the common logistics principles.

The information system “Logistics of the BA” (IS ‘Logistics’ of BA) is built in accordance with the Conception for Information Strategy of MOD referring to C4I systems ensuring the necessary compatibility with the systems of the other countries members of NATO. The project aims at fulfilling the requirements for building a modern information system, which will ensure: [4]

- improving the efficiency of management of logistics support
- increasing opportunities to integrate with the other NATO members armed forces
- providing actual info at actual time for material and financial resources in BA
- providing means for control and management of the processes in BA and MOD in real time

The system is designed to automate and optimize the work of the logistical structures, by providing timely and accurate reference and management information. By way of its component parts, it can be referred to as ERP last generation system.

Initially, the system was programmed to serve the logistical needs of the BA, but it has been provisioned for it to be able to widen its scope to include also MOD with regards to its structures for finance, planning, budgeting, and material resources procurement.

IS Logistics of BA is a web based system for monitoring, management, and control of the logistical and financial processes, where the customers find a centralized data base in servers, which save all the info for all the BA units. This info is not available at the regular working place of the customers or in the units.

The access of the users to the resources of the system is restricted. It is in close dependence to the specific functional responsibilities of the logistical bodies. At the same time, the system is entirely transparent and doesn't allow abuse, or fraudulent actions on the part of the users. As soon as it starts an operation with material resources, it becomes visible to the users who are allowed access to it.

IS Logistics of BA is an integrated system. It is a complex and an instrument to manage and execute all the ordered movement with the material resources, including automated registering of all the accounts sheets in a previously set accounts module.

This system uses Oracle Applications and Oracle Database, adapted for the military. In the different functional modules of these software products many world recognized practices are applied (logistical, accounting, planning, contest announcing, collecting offers, conducting auctions and others) and the ways these are carried out.

The development and implementation of *IS Logistics* of BA starts in 2003 with a pilot project. It is funded by the Program for foreign financing in the military (FMF). An agreement memorandum and subsequently letters of intentions and requirements to the system have been signed between the MOD and the Office of Military Cooperation of the American Embassy in Bulgaria.

The executive person of the project is the American company Unisys Federal Corporation, with sub-executives the Bulgarian firms Solutions Ltd.- one of the leading firms, implementing solutions on the basis of ERP software from the world recognized business software suppliers- Oracle and Microsoft. The project has been carried out under the surveillance of TSIO (Theater Systems for Information Office) attached to NATO, and together with Manifold Corporation.

The execution of the pilot project will go through three stages:

Stage1: (2003-2005). During this stage, the main functional areas and work places in the units are developed, which were initially included in the Pilot Project of the Logistics Information system. A trial of the system has also been conducted under limited organizational and functional scope.

Stage 2: (2007-2009). During this period, the system software has been further developed, as the functions have been increased to include the greater part of the army logistical procedures. Users' feedback tests have been applied, which confirm the functional range and its correspondence to the applied document papers, as well as its performance. A six-month long laboratory test examinations have been carried out to find and remove faults in the system.

Stage 3: (2012-2014). This is the final stage to implement the information system in the military units of the BA and the MOD structures included in the project. Additionally, during this period the software is being improved, by gradually moving from 32 to 64 Bits version, the infrastructure is being built(communication points and work stations) in 10 more military units that will be included in the system. Also training of the personnel that will work with the system is being carried out. Test for the completion stage are also done, a final operation test of *IS Logistics* of BA, included in that number.

According to the signed memorandum, there is a provision for 12 months software maintenance by the person in charge of the execution of the project, and as far as hardware is concerned, the maintenance will be with respect to the warranty guarantees of the producer companies.

IS Logistics of BA presents a mechanism and ready-made procedures and algorithms of operation and management of more than 200 logistic and finance processes. The separate modules and the logistical processes are closely related, and work in a synchronized way while sharing the data needed. At every stage of the system operation, there are processes of managing the finances and managing the material resources reserves.

On a functional level the system comprises unified standard procedures and operational processes in the following spheres: planning, supply and delivery, storage and preservation; application and usage; management of the material resources; technical support; upkeep; individual object safekeeping; finance and accounts management.

Module Planning ensures the distribution of the financial means among the organizations, registered in the system, the development of a unified plan for material and technical support (UPMTS) and the drawing of the annual budget for each organization. The operational modes for

the completion of a number of activities are included in the planning module: statements for financial limits for programs or paragraphs, planning of material resources, munitions limits, usage of railway transport limits, limits for the disposal of the redundant material means and utilization, for the individual object safekeeping of the personnel and so on.

Module Supply controls the supplies made to the system from internal and external sources. By means of the logistics processes in this module, the accomplishment of the contract provisions and receiving deliveries is observed and controlled- as far as quality and quantity is concerned, as well.

Modules Storage premises and Fixed assets manage the stock and safe storage of MS (Material supplies) and FA (fixed assets) of the organizations, registered in the system.

Module Storage allows for a flexible structure of storehouses to be built, with its subdivisions of storehouses and location devices for them. The warehouses can be set to give balance and off-balance accounts sheet, and the average estimated cost of each material stock, is calculated automatically.

These modules cover all the working processes of MR and FA- registering the in-stock and out-stock, developing an inventory of assets, re-arrangement, culling and others. The accounts sheets are set for each separate logistical process and the system automatically registers the accounting data, produced by the material operations.

Module Material Resources Usage allows that all the resources used are monitored and registered with view to their target use as well./ to be spent or to be returned/. The systems makes it possible to work with practical an indefinite number of materially-responsible personnel, with the use of locators for that effect.

Module Finance ensures budget control and automatic accounts record of the financial and economic operations done in the system. This module works in close connection with the Planning Module, as it ensures the control and accounts statements of the financial resources respectively for the programs, paragraphs, and positions of common material plan of Bulgarian Army.

The information system has its own accounting modules, corresponding to the national and international accounting standards and assigns a unified accounts plan for all organizations, registered in the system. The accounts record data and the required reference data for each separate organization are available for control and examination by the higher ranking authorities at real times. The system keeps at its disposal, a wide number of reference and accounts record sheets, and some of these, on request, can be consolidated.

Also the system keeps a big number reference and accounts sheets used by the logistics structures in their work and they cover all the functions in the system.

IS Logistics of BA is founded on an hierarchical basis, based on its turn, on the chain of command of the BA. The system provides means for hierarchical units to be built, according to the set objectives. These units are built dynamically and can be altered with time.

In their scope a number of outside BA structures can be involved, if their activities are realized through the

logistics system. At present these functions are realized through structures in the system.

This hierarchical structure also lies in the completed model for security of the system. The customers are allowed access according to their clearance and this access is in relation only to their own respective organizational units and sub-units.

The scope of the model depends on the scope of the pilot project, the customers, and the requirements for the cooperation of the system with outside systems. Defining the particular levels of interaction with the system through users interface, we can come down to the following:

- Participants- organizational units (departments, units, command structures), interacting with the system.

- Work places- there will be a few work-place in each organizational unit, each with a range of defined separate functions to which there is an access.

- Users- any work place, but it is not recommendable to be used by several users, since each one has a defined range of functions they have access to. One user can be authorized to deal with all the functions for this working place, or a part of them.

The requirements for information access for the different levels of the armed forces structures, are also different with respect to the range of data access and functions performed. The data access is based on the principle that from each organizational unit, the subdivisions data record can be tracked, down the hierarchical chain. The high level users in charge of one type of MR (material resources), can track down those particular items of MR down the chain.

The system users are the officials, authorized to work with *IS Logistics* of BA.

From the users point of view the interaction with the system depends on what report data goes in, and what results come out. Users of the systems will be: commander; accountant; MR officer; SPC Technical support; SPC Military Assets; materially responsible officer/MRO/ of the military store, Application Administrator; System security Administrator and Database Administrator.

Towards the end of 2013, the basic tasks of Stage 3 were completed and trial implementation tests of the system were done, which tested its: functionality, the configuration of the system, developed benchmarks, the system operability. The system operation performance was checked with different numbers of users and a particular number of trial situations in operation mode.

A test, establishing a standard, was conducted. It tests the time needed for the system to respond to queries from different locations of the transfer area. Also a final operative test was done, testing the system in real production environment, through manual execution of a certain number of scenarios.

Confirmation trials show that after some improvements to the software have been done, it works properly and ensures the automated logistics management.

IS Logistics of BA will considerably improve the logistical work in the Bulgarian army and MoD, in all aspects, and will allow management decisions of every character and level to be made- strategic, operative, tactical. The existence of such a system, is also one of

NATO requirements for the further development of our armed forces.

The system will provide timely information for the in stock material resources, needed by the army. It will be useful for the management personnel on strategic and operative levels.

The realization of a project of such a scope holds a lot of difficulties, though. Many people do not grasp its importance, and also do not understand the system in its essence and objectives. It's very hard to part with old habits and work methods, established in the past decades. The problem with the lack of well-trained staff for the realization of the project is also an issue. The motivation of the personnel is not up to what it should be, mainly, due to the fact that together with the work on the implementation of the system, they are to manage the logistical functions for the army, as well.

There is no other alternative, but to develop the *IS Logistics* of BA. The opposite would mean that instead of going with the flow in the era of super fast speed and information exchange, we have chosen to go back to the stone age. At present, everything comes down to persistence, hard work, and the governing will to train well-prepared and motivated staff. A lot of hard work should be done, in order to accomplish the assigned tasks.

III. ANALYSIS OF THE BENEFITS OF THE APPLICATION IS LOGISTICS OF BA FOR THE STORAGE IN BA.

The expected benefits of the systems are great, since the investment put is considerable. While defining the systems' characteristics, some of these benefits have already been outlined; such as real time information, allowing for better management decisions to be made, and paperless exchange of information.

As a result of the research made, certain conclusions referring to its benefits can be made: [5]

- Improvement of the army organization and management related to materials' supply and services provided
- Optimized reserves management, leading to prevention of surplus storage of stock
- Improved organization and accountability of storehouse goods- in all areas- provision, equipment, munitions and so on.
- Opportunity for better planning and better realization of supply refreshment
- Information assisted auctions and contract agreements
- Optimal finance control of operations with materials and services.
- Better options for the command staff to control the storehouse access.
- Unified names for materials and services in BA and their compatibility to other coding systems, for instance NATO Stock Number
- Providing ways for data transfer with other NATO systems.

These benefits, so outlined, apply to the whole logistics system in the BA. It will be useful, however, to examine also the concrete benefits for the warehousing in BA- specifically in the spheres of: introduction of the Bar Code technology, and also in relation to the amount of data

reports pertaining to the basic and additional parameters of the material resources- the object of safe-keeping. [6]

First, the barcode technology is to be introduced.

This is a technology that has been in use for 58 years. It was originally patented on 7th Oct., 1952 after the names of Joseph Woodland and Bernard Silver. Initially the barcodes were used for car labeling, but they didn't get universal recognition until the grocery stores and later supermarkets, started to use them, which drastically speeded the service to the customers. Barcode technology has been constantly developing and today we can find different systems for coding with different symbols. Their mass utilization helps considerably the business processes connected to the tracking of stock, sales, sending and tracking packages etc.

The usage of barcodes is provisioned in IS Logistics of BA, only for some processes, namely: barcodes with continuous access and bar codes with non-continuous access. All MR(material resources) in the project input, are an object of identification according to EAN (European Article Numbering) standard and have to have identity numbers. Trade item are identified by EAN number / supplier, client, carrier, firm/, logistic units/containers / and locations. Each trade item is coded by GTN/ global trade number/, which has 14-digits code structure, which allows us to distinguish between single items, packets, and container/pallets, that is, it covers all levels of packaging. The code scan gives information for the features in several data sheets. Such as, for example, in the software the system uses, several data will be filled in: item name, shipment number, quantity and measurement figures.

With applying **the continuous access barcodes**, the software makes it possible for the processes to be managed by means of bar code devices. To use them, some conditions should be fulfilled: the material object should be registered by its barcode in the system, and there should also be a connection between the manufacture barcode and the MR in the nomenclature storehouse. To work with barcodes we need to build a communication structure. For this purpose a network should be developed, with storehouse scope, and also to buy particular barcode devices, combined with barcode printer.

The processes where this new technology can be applied are as follows:

- In-put of MR with barcode read
- Check for ramp correspondence
- Management of containers and ID labeling
- Allocation of subdivision storehouses- zones, coding of locators
- Partial and complete inventory
- Tracking of validity and expiry terms of MR
- Packaging and parceling
- Preparation, assembling and loading of shipment and others.

Using the non-continuous access barcodes, where there is no constant connection between the barcode read and the info system, it's necessary that the devices and the system, have a synchronization protocol. When a connection is established the system can work efficiently, because it receives the data needed. Such data are: a list of nomenclature for the particular storehouse, the ID of the

user of device, list of in-stock, as well as other information. After the device has been initially charged, it can be disconnected from the system and it goes to autonomous work mode. This makes it possible for the materially accountable person to carry out a number of warehouse activities, as the results are logged in the specially designed data base. The work finished in the autonomous mode, the materially accountable person connects the device to the centralized system again. It can be done out of the warehouse perimeter. After the synchronization, all the data in the mobile device are deleted, and it is ready, once again, to be initially charged.

Functional uses of barcodes with non-continuous access:

- Receiving stock from supplier
- Shipping stock to other warehouses in the system scope/ internal shipping/
- Automatic ID of labels
- Location control- with the creation of the necessary organization of the warehouses, the system will make it possible for data in-put, identifying the items' location in warehouse, by the use of barcode labels.
- Partial and complete inventory and others

In the second place, the information system can ensure for the storehouse management benefits in relation to the amount of information kept referring to basic and additional parameters of the MR- the object of safe-keeping.

Logistics information systems use different parameters to identify the stock and services through their entire life cycle, and they also differ in the processes used for identification and demanding for them- planning, purchase, delivery order, storage and preservation, exploitation, until they are waived of report.

The identification and specification of the basic and additional parameters of the material resources and services, together with the respective business processes, is a key point in the development of any logistic information system, as this classification, although theoretical, gives a clear idea of the extent of use of those parameters in a given functional environment and the processes involved.

It's necessary to distinguish between the following parameters: Type MR; Nomenclature N in BA; Name; Declared manufacturer; Nomenclature number of manufacturer; Substituting stock; Measurement figures; Quantity; Category of MR; Cost.

The parameters of the material resources should be viewed with respect to the desired functionality of the information system, and in relation to their connection to the NATO systems for classification and coding, namely:

- in connection with the processes of planning and management of deliveries: packaging type; physical proportions; weight; place of delivery; time of delivery.

- in connection with the processes of storing and preservation of material resources: delivery date; manufacturer/supplier; manufacture date; validity period; shipment number; serial number; conditions for storage and preservation, assigned by manufacturer, period of preservation at storage, assigned by the manufacturer, period of preservation at field conditions, assigned by manufacturer, storehouse/place of preservation; coordinates in store; actual conditions of preservation/

type conservation; date of preservation/conservation; period of preservation/conservation.

- in connection to the processes of exploitation: conditions of exploitation; maximal technical resource; period of technical servicing; type of technical servicing; actual conditions of exploitation; processed resources; technical condition; technical resource until next technical servicing.

IV. CONCLUSIONS

On the basis of the research done, we can conclude:

- The application of IS Logistics of BA holds a number of **hazards**, too; insufficient amount of nomenclature material stock in data base; not enough data about parameters and features of MR; out-of-date data base; incorrectly assigned users of the system- necessary officials not involved; key users not determined; commanders not involved; officials responsible not involved; no provision and delays in the technical implementation of the system; insufficiently trained users; interruptions of project development due to financial reasons.

- **Main benefits** of the implementation of the system for the Bulgarian armed forces: improvement of organization and management of supply; optimal reserves control; improvement of store-house organization and accounts reports; opportunity for better planning and realization of the refreshment of reserves; information assisted auctions and contract talks; full financial control of operations; better command control on the warehouse access; summed and detailed information about all logistics and financial operations; unification of names of MR and services and correspondence to foreign coding systems; ensured data share with other NATO systems.

- **The concrete benefits** for the storehouse management of BA are: the use of the barcode technology; the upkeep of full data base about the MR parameters at preservation; real time information exchange; generated information flows for the needs of control and management.

We have to conclude here, that the establishment and structuring of a data base, the accurate determining of the users of the system, as well as, the timely and accurate information for the supplies in stock and the demands, for the factual progress of material flows, can reveal the whole potential of the IS Logistics of BA, and the real benefits it brings for the warehouse management of the BA.

The matter of how useful it is, will yet to be confirmed and reconfirmed in the years to come, when it is going to still bring more new benefits for the management of the warehouses of the Bulgarian army.

REFERENCES

- [1] Concept for the development of the Bulgarian armed forces, MoD, S., 2010, p. 23-24.
- [2] Nichev, N. *Theoretical foundations of military logistics- organization of logistical support*. Economics Dept., at Vassil Levski NMU, V.Tarnovo, 2011., p.10
- [3] Report about the state of defense and armed force during 2010, MoD, S., 2010, p. 44.

- [4] http://cio.bg/6508_logistichnata_is_na_ba__proekt_bez_alternativa - 18. 05. 2015.
- [5] Ivanov, Iv. Directions for improvement of Bulgarian army's warehouses. *G. S. Rakovski, Military Academy*, S. 2006, p. 256.
- [6] Stoianov, M. Information system Bulgarian Army Logistics – essence, development and priorities in her application in material storage. *at Vassil Levski NMU*, V.Tarnovo, 2011, p. 95-100.

Further ‘in silico’ validation of a facial affect recognition system

Peter Vida^{1,2}, Jozsef Halasz^{2,3}

1. School of Doctorate Studies, Semmelweis University, Hungary

2. Vadaskert Child Psychiatry Hospital, Budapest, Hungary

3. Obuda University, AMK, Szekesfehervar, Hungary

fitz026@yahoo.co.uk

halasz.jozsef@amk.uni-obuda.hu

Abstract— The recognition of human facial affect has crucial importance in social communication and proper social functioning. In certain psychological and psychiatric conditions, the recognition of facial affect and the overall social functioning are also impaired. In recent years, social prosthetics occurred, giving a possibility for later therapeutic applications in these conditions. In the present paper, further characterization of a facial affect recognition system, FaceReader is described, focusing on three basic emotions: anger, fear, and happiness. With the usage of controlled expression intensity and laterality, present data indicate that the above system is capable for discriminating basic emotions and is sensitive for stimulus intensity. In the case of anger, no major effect of laterality was present, but the system was sensitive for laterality in the case of fear and happiness.

Keywords: Affect Recognition System, Anger, Facial, Fear, Happiness, Social.

I. INTRODUCTION

In the late sixties of the last century, six basic universal emotions were described (anger, disgust, fear, happiness, sadness and surprise) as pan-cultural elements of facial affect [1]. The perception and interpretation of facial affect in humans is a complex multilayer mechanism for each emotion [2]. A growing interest for facial affect recognition systems occurred in the field of education, market research, psychology and computer sciences [3-5], but the usage of these systems as “social prosthetics” in vulnerable individuals was also suggested [6]. The discrimination of these basic facial emotions is markedly important in the social functioning of the individual. Impairment in the recognition and the interpretation of these patterns might accompany different psychiatric conditions, among them autism spectrum disorder, schizophrenia and antisocial development [7,8]. So far it is unclear whether specific recognition deficits might occur in specific psychiatric conditions, but the usage of these systems might have major potential diagnostic and therapeutic effect in the above psychiatric conditions [6]. Our original research line targets adolescents with externalization problems, and this population is vulnerable towards antisocial development. A sensitive issue within this population is the marked difficulty in the recognition of fear [9-11]. This specific impairment might result

controversies in dyadic encounters, e.g. the misinterpretation of the occurring fear on the face of the opponent might result the continuation of a physical fight in adolescents. Targeted therapeutic inventions were initiated to decrease the level of “fear blindness” in adolescents with externalization problems [12]. Using validated systems capable for automatic detection of facial affect in adolescents with conduct problems might also outline major alterations responsible for behavioral patterns. Thus, for the above reason, only facial emotion recognition systems with established and detailed characteristics can be used in clinical studies.

The Noldus FaceReader system was designed for fast and reliable automated analysis of facial expressions [13]. The external validation for static inputs on a large number of input images was performed earlier [14,15], and the system was able to discriminate effectively according to the six basic emotions. In the above papers, only the first-order emotion matrices were described. Later publications on validated static inputs also confirmed the usability of the system [16,17], and in a recent paper the detailed response field for validated static inputs was also described [18]. Still, certain aspects of the system were not described, what might be important for later interpretation of using the system in patients with different psychiatric conditions. Validated databases were sensitive to the nature of emotion, but the intensity of the emotions was not clarified. Additionally, the above papers used images with faces in frontal position, thus to our best knowledge, neither the intensity nor the laterality of the visible faces were controlled in the above studies. Basically, a new set of pictures with controlled intensity and laterality might be appropriate for the above trial. Instead of these controlled experiments, a set of ‘in silico’ faces were created to test the effect of laterality and intensity in the Noldus FaceReader system.

The aim of the present study was to describe the emotion specific responses onto stimuli with controlled intensity and laterality on three different emotions: anger, fear and happiness. A marked effect of expressed emotion intensity was suggested, while no effect was suggested in the case of slight laterality. In our previous study, the system was less sensitive in the recognition of anger, and the most sensitive in the recognition of happiness [18]. The recognition of fear has a major importance in antisocial development [9-11], as described above, thus we also studied the detection profile on expressed fear.

II. METHODS

In the present set of experiments, facial masks were generated (with different expressed emotions) in the free version of the software FaceGen Modeller 3.5 (Singular Inversions). Sample images were presented in Fig. 1 and Fig. 2, respectively. In the first set of the images, a frontal view of the images was present (Fig. 1), while a standardized lateral view was applied in the second set of images, with a moderate laterality of 20 degrees (Fig. 2).

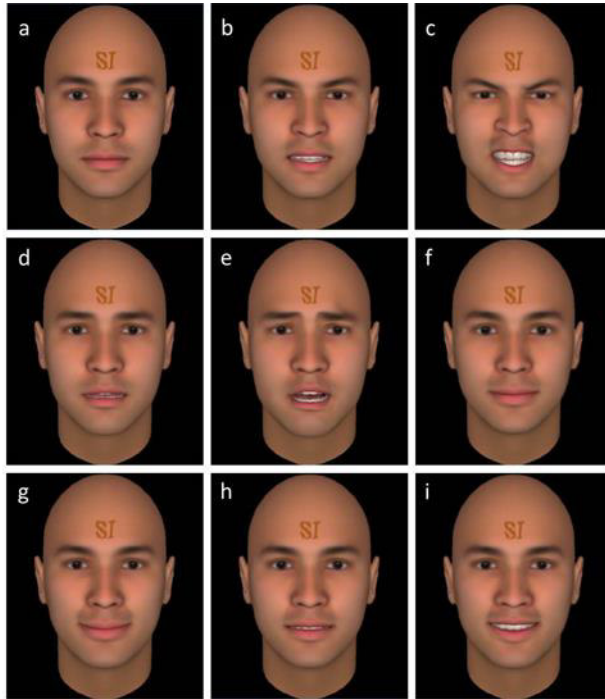


Fig. 1. Sample frontal artificial images from the FaceGen Modeller; a, neutral; b-c, anger low and high intensity; d-e, fear low and high intensity; f-g, happiness (mouth closed), h-i, happiness (mouth open) low and high intensity.

The same set of artificial masks was applied in both camera positions. This procedure was selected in order to establish full controllability and standardization of the images. Additional to the images expressing neutral emotional content, the following emotion intensities were modulated: anger, fear, happiness. As it was described earlier, the present paper was based on earlier observation with validated, standardized images. From the six basic emotions, images expressing anger, fear and happiness were selected. With the FaceGen Modeller, the emotional content was selectively modified in 0.1 steps (10 percentage) between neutral values and maximal expression of the given emotions. In the case of happiness, both open and closed mouth versions were tested. The average responses for 0.1-0.5 intensities were considered as low intensities, while 0.6-1.0 emotional intensities were considered as high intensities. Altogether, 110 images were analyzed by the means of the FaceReader 5.0 program. The detailed emotional responses for a given image were normalized, and neutral

responses were also calculated [11]. In the present paper, emotion-specific responses were compared.

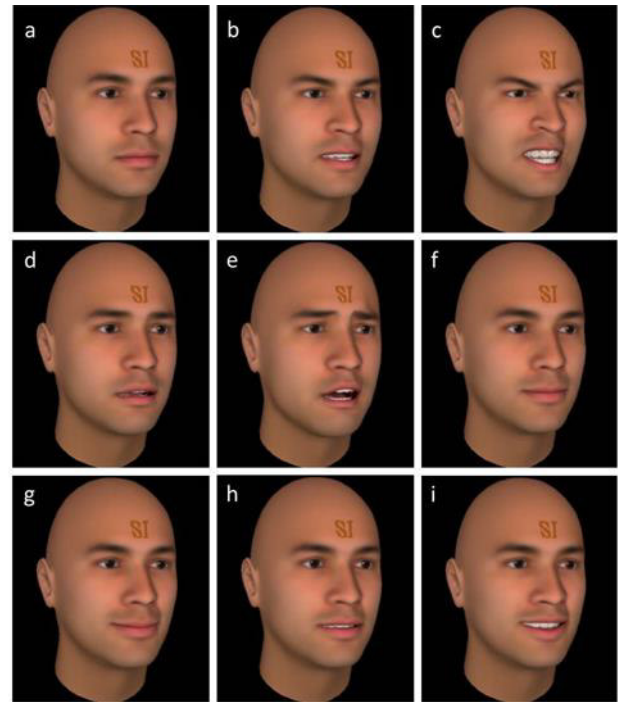


Fig. 2. Sample lateral artificial images from the FaceGen Modeller; a, neutral; b-c, anger low and high intensity; d-e, fear low and high intensity; f-g, happiness (mouth closed), h-i, happiness (mouth open) low and high intensity.

Statistical analysis. Statistica 7.0 program was used for the analysis, GLM model was applied. Data are expressed as means and SEM (Standard Error of the Mean). Laterality (frontal vs. lateral) was used as repeated variable, while the intensity of affect (low vs. high) and the interaction between laterality and intensity was also calculated. Newman-Keuls post hoc comparisons were run where appropriate, $p < 0.05$ was considered as significant difference between groups.

III. RESULTS

In the case of anger stimuli, a significant effect of emotion intensity responses were present ($F_{(2,14)}=6.670$, $p < 0.01$), while no effect of laterality was observed ($F_{(1,14)}=3.138$, $p < 0.1$). Interestingly, higher intensity on the expressed masks was accompanied by higher detection responses only in the case of lateral view (Fig. 3).

In the case of fear stimuli, a significant effect of emotion intensity responses were present ($F_{(2,14)}=10.980$, $p < 0.0014$), and marked effect of laterality was also present ($F_{(1,14)}=11.769$, $p < 0.005$). Low fear intensities did not differ from the baseline. High fear intensities were properly discriminated in the frontal view (Fig. 4). In the case of lateral view, a marked impairment was visible in the detection.

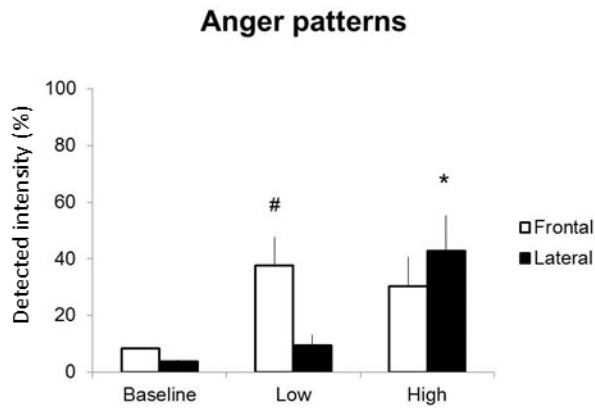


Fig. 3. Anger patterns. Data are expressed as means±SEM. Low, stimulus with low intensity; High, stimulus with high intensity; Frontal, frontal artificial images; Lateral, lateral view of images (20 degrees); *, significantly different ($p<0.05$) from baseline; #, significantly different ($p<0.05$) from frontal view.

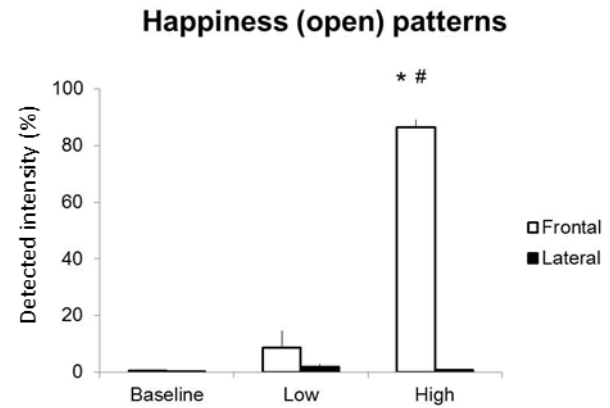


Fig. 6. Happiness (mouth open) patterns. Data are expressed as means±SEM. Low, stimulus with low intensity; High, stimulus with high intensity; Frontal, frontal artificial images; Lateral, lateral view of images (20 degrees); *, significantly different ($p<0.05$) from baseline; #, significantly different ($p<0.05$) from frontal view.

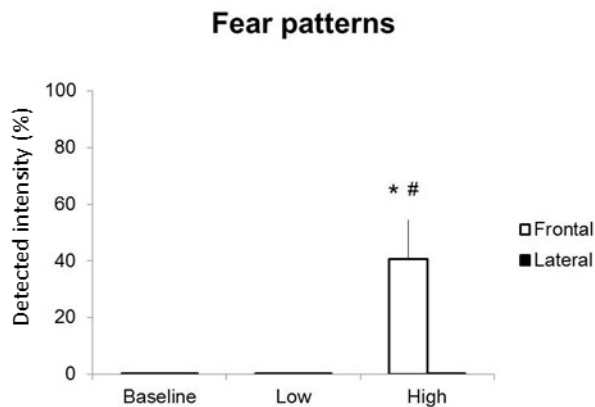


Fig. 4. Fear patterns. Data are expressed as means±SEM. Low, stimulus with low intensity; High, stimulus with high intensity; Frontal, frontal artificial images; Lateral, lateral view of images (20 degrees); *, significantly different ($p<0.05$) from baseline; #, significantly different ($p<0.05$) from frontal view.

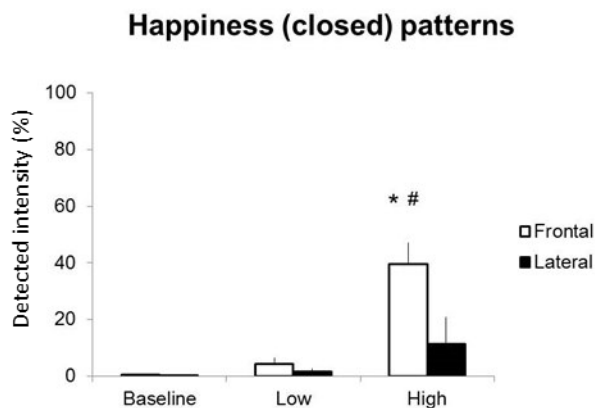


Fig. 5. Happiness (mouth closed) patterns. Data are expressed as means±SEM. Low, stimulus with low intensity; High, stimulus with high intensity; Frontal, frontal artificial images; Lateral, lateral view of images (20 degrees); *, significantly different ($p<0.05$) from baseline; #, significantly different ($p<0.05$) from frontal view.

In the case of happiness stimuli, both the effect of closed and open smile was analyzed. In the case of happiness with the mouth closed, a significant effect of emotion intensity responses were present ($F_{(2,14)}=13.994$, $p<0.001$), and marked effect of laterality was also present ($F_{(1,14)}=12.547$, $p<0.004$). Low intensities did not differ from the baseline. High closed happiness intensities were properly discriminated in the frontal view (Fig. 5). In the case of lateral view, a marked impairment was visible in the detection. In the case of happiness with the mouth open, a significant effect of emotion intensity responses were present ($F_{(2,14)}=190.356$, $p<0.00001$), and marked effect of laterality was also present ($F_{(1,14)}=262.603$, $p<0.00001$). Low intensities did not differ from the baseline. High open happiness intensities were properly discriminated in the frontal view (Fig. 6). In the case of lateral view, a marked impairment was visible in the detection.

IV. DISCUSSION

The main results of the present study were the following. First, the system was most sensitive for the detection of 'in silico' happiness, especially when the mouth was open during expression. Second, in the case of fear and happiness, the FaceReader was sensitive to the intensity of the stimuli, and the slight lateral view significantly impaired the performance. Third, in the case of anger, no effect of laterality was present, and low intensities resulted high responses in the frontal view.

According to earlier reports, FaceReader has the highest recognition sensitivity in the case of happiness among the basic emotions [14-18]. These studies did not discriminate between closed and open mouth patterns. Our data indicate that happiness with open mouth is a stronger signal to the system, but reliable data can be acquired even with happiness with mouth closed. The closed pattern might also be considered as a state with lower intensity, thus the data outline the sensitivity of the system for different intensities.

In the case of fear, only high intensities were recognized, and the recognition pattern was similar to pattern visible in the case of closed happiness. In humans, among the basic emotions, the most difficult is the detection of fear [19]. Albeit the FaceReader 5.0 system was sensitive to the intensities, slight laterality abolished the effect.

Laterality did not cause problem in the case of anger, albeit slightly modified the detection pattern seen in the case of frontal view. This is a differential effect, and can be taken into account for later analysis. E.g., when someone rotating the head during the experiment, the proportion of angry responses might relatively get increased compared to the proportion of other responses. In the case of externalizing adolescents, where comorbid attention deficit hyperactivity (ADHD) might be present, the above effect might be taken seriously [20].

The limitation of our study was that ‘in silico’ faces were used for testing the effect of stimulus intensity and laterality. Further studies might include validated databases (both in terms of emotion specificity, intensity and laterality), but at present, we have no information on such a database. On the other hand, this computerized method was able to establish standardized measures in terms of angle and intensity, which is definitely robust (if possible) effort from a human presenter.

In conclusion, these data suggest that stimulus intensity might be detected with the FaceReader 5.0 system, and happiness is the most sensitive among the emotions studied. Frontal view is necessary to acquire the above result, and the effect of even slight laterality must be taken into account.

ACKNOWLEDGMENT

The author declares no conflict of interest. The work was supported by OTKA Grant No. 108336.

REFERENCES

- [1] P. Ekman, E.R. Sorenson and W.V. Friesen, “Pan-cultural elements in facial displays of emotion”, *Science*, vol. 164, pp. 86–88, 1969.
- [2] R. Adolphs, “Recognizing emotions from facial expressions: psychological and neurological mechanisms”, *Behav. Cogn. Neurosci. Rev.*, vol. 1, pp. 21–62, 2002.
- [3] J.M. Harley, F. Bouchet, M.S. Hussain, R. Azevedo and R. Calvo, “A multi-componential analysis of emotions during complex learning with an intelligent multi-agent system”, *Comp.Hum. Behav.*, vol. 48, pp. 615–625, 2015.
- [4] P. Lewinski, M.L. Fransen and E.S.H. Tan, “Predicting advertising effectiveness by facial expressions in response to amusing persuasive stimuli”, *J. Neurosci. Psychol. Econ.*, vol. 7, pp. 1–14, 2014.
- [5] M. Gadea, M. Alino, R. Espert and A. Salvador, “Deceit and facial expression in children: the enabling role of the “poker face” child and the dependent personality of the detector”, *Front. Psychol.*, vol. 6, Article 1089, pp. 1–9, 2015.
- [6] J. Halasz, N. Aspan, P. Vida and J. Gadoros, “Recognition of emotions in facial expressions- implications and limitations during antisocial development”, 7th International Symposium on Applied Informatics and Related Areas, pp. 44–49, 2012.
- [7] C. Herba and M. Phillips, “Annotation: Development of facial expression recognition from childhood to adolescence: behavioural and neurological perspectives”, *J. Child. Psychol. Psychiatr.*, vol. 45, pp. 1185–1198, 2004.
- [8] L. Collin, J. Bindra, M. Raju, C. Gillberg, and H. Minnis, “Facial emotion recognition in child psychiatry: a systematic review”, *Res. Dev. Disabil.*, vol. 34, pp. 1505–1520, 2013.
- [9] M.R. Dadds, Y. El Masry, S. Wimalaweera and A.J. Guastella, “Reduced eye gaze explains “fear blindness” in childhood psychopathic traits”, *J. Am. Acad. Child Adolesc. Psychiatry*, vol. 47, pp. 455–463, 2008.
- [10] N. Aspan, P. Vida, J. Gadoros, and J. Halasz, “Conduct symptoms and emotion recognition in adolescent boys with externalization problems”, *Scientific World Journal*, vol. 2013, p. 826108, 2013.
- [11] J. Halasz, N. Aspan, C. Bozsik, J. Gadoros, and J. Inantsy-Pap, “The relationship between conduct symptoms and the recognition of emotions in non-clinical adolescents”, *Psychiatr. Hung.*, vol. 28, pp. 104–110, 2013.
- [12] M.R. Dadds, A.J. Cauchi, S. Wimalaweera, D.J. Hawes and J. Brennan, “Outcomes, moderators and mediators of emphatic-emotion recognition training for complex conduct problems in childhood”, *Psychiatr. Res.*, vol. 199, pp. 201–207, 2012.
- [13] Noldus Information Technology B.V., “FaceReader 5.0 Reference Manual”, Noldus GMB, The Netherlands, pp. 1–154, 2012.
- [14] V. Terzis, C.N. Moridis and A.A. Economides, “Measuring instant emotions during a self-assessment test: the use of FaceReader”, *Proceed. Measur. Behav.*, pp. 192–195, 2010.
- [15] V. Terzis, C.N. Moridis and A.A. Economides, “Measuring instant emotions based on facial expressions during computer-based assessment”, *Pers. Ubiquit. Comput.*, vol. 17, pp. 43–52, 2013.
- [16] P. Lewinski, T.M. Den Uyl and P. Butler, “Automated facial coding: Validation of basic emotions and FACS AUs in facereader”, *J. Neurosci. Psychol. Econ.*, vol. 7, pp. 227–236, 2014.
- [17] M. Olszanowski, G. Pochwatko, K. Kuklinski, Scibor-Rylski, P. Lewinski and R.K. Ohme, “Warsaw set of emotional facial expression pictures: a validation study of facial display photographs”, *Front. Psychol.*, vol. 5, Article 1516, pp. 1–8, 2015.
- [18] J. Halasz, N. Aspan, P. Vida and J. Balazs, “Output properties for validated static inputs in a facial affect recognition system”, 9th International Symposium on Applied Informatics and Related Areas, pp. 25–30, 2014.
- [19] A.A. Marsh and R.J. Blair, “Deficits in facial affect recognition among antisocial populations: a meta-analysis”, *Neurosci. Biobehav. Rev.*, vol. 32, pp. 454–465, 2008.
- [20] National Institute for Health and Care Excellence, “Antisocial behaviour and conduct disorder in children and young people. Recognition, intervention and management”, British Psychological Society and The Royal College of Psychiatrists, *National Clinical Guideline Number 158*, pp. 1–468, 2013.

Application of modern computer methods for acoustic phenomena investigation

T. Trifonov*, I.S. Ivanov**

*St. Cyril and St. Methodius University of Veliko Tarnovo /Department of Computer Systems and Technologies and National Military University "Vasil Levski"/Department of Communication and Information Systems, Veliko Tarnovo, Bulgaria,

**National Military University "Vasil Levski"/Department of Communication and Information Systems, Veliko Tarnovo, Bulgaria

e-mail: tihomirtrifonov@ieee.org, ivanov_ivan@nvu.bg

Abstract—In this paper a modern method of investigating and acoustic analyzing of some natural phenomena in urban environment is presented. Comparison between sound of church bells and battle systems is carried out. It is shown, that the infrasound component has very strong influence on noise energy. Wavelet analysis is proposed for signal processing. The opportunity of MatLab for modeling and calculation of applications is used.

I. INTRODUCTION

A few years ago it was thought that in the field of mechanical vibrations and acoustical phenomena everything is studied and everything is known. Recently, however, with the development of new measuring instruments and new methods of mathematical analysis, it became clear that new opportunities appear for in-depth study. Due to the limited possibilities of classical technologies it was not possible in the past. It is interesting that some of these opportunities are directly connected with the computer data preservation and protection, [1].

II. MODERN ASPECTS OF HARMONIC ANALYSIS

When we use the "harmonic" throughout mathematics (harmonic analysis etc.), we make the link with music, because its origin lies in the music theory.

Pythagoras (582-507 BC) provided mathematical tools that would prove very useful in the calculation in many fields (music, astronomy). But today, historians know that Pythagoras and his followers were not the first to make use of the famous theorem for a right triangle. There is some evidence that the Egyptians made use of it for many centuries before the Pythagoreans.

Milesian philosophers including Pythagoras's teachers, Thales and Anaximander, were familiar with Babylonian science, [2].

In modern times Clairaut, Lagrange, Daniel Bernoulli, Euler and Gauss developed the ideas of a harmonic analysis, and Joseph Fourier made the breakthrough development with Fourier series in your 1807 paper "*Mémoire sur la propagation de la chaleur dans les corps solides*".

Today, the fast Fourier transform-FFT algorithm is one famous realization of these ideas. FFTs were first discussed by Cooley and Tukey (1965), although it may

be noted that Gauss in one of his work (dated to 1805) shows the use of a particular technique of separation, which forms the basis of modern FFT algorithms. Cornelius Lanczos did pioneering work along with Danielson about FFT, 1940 but the significance of his discovery was not appreciated at the time.

In last decades new ideas are being applied in harmonic analysis – wavelets, that have some advantages over traditional Fourier analysis. Starting with Haar's functions, Gabor atoms, and today Mallat and Daubechies and other families of wavelets [3] wavelet analysis become very useful tool in advanced digital signal processing.

III. CONFORMAL MAPPING OF SCALOGRAM

Continuous wavelet transform (CWT) gives the improvement for analysis via conversion of 2D signal into pseudo 3D signal - scalogram.

The usual scalogram has the form of a strip. After appropriate changes of the position and its dimensions we transform the scalogram by the exponential function onto a concentric ring. In fact, this conformal transformation gives us some advantages in analysis, [4].

Let $w = f(z)$ be a complex valued function of complex variables. If $f'(z_0) \neq 0$, then

(i) in some neighborhood of z_0 the function $f(z)$ is univalent, i.e.

$$f(z_1) = f(z_2) \Rightarrow z_1 = z_2;$$

(ii) the function preserves the angles between curves intersecting in the point z_0 , i.e. there the mapping realized by $f(z)$ is conformal.

The exponential function $w = e^z$ is the most interesting function from analytic and geometric point of view:

(i) e^z is defined and different from zero for each $z \neq \infty$;

(ii) $(e^z)' = e^z \neq 0$ for each $z \neq \infty$; (cf. above remark for univalence);

(iii) $e^z = e^{x+iy} = e^x \cdot e^{iy}$ shows that:

(a) lines parallel to the abscise, mapped by e^z become rays starting from the origin;

(b) lines parallel to ordinate axis, mapped by e^z become circles centered at the origin;

(c) the exponential function is periodic with basic period $2\pi i$.

The last property says that the exponential function is univalent, in the above described sense, in each horizontal strip with (vertical) width 2π .

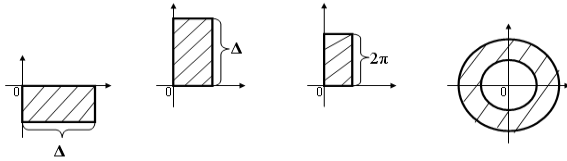


Figure 1. Conformal transform of scalogram (rectangular)

Let $-\infty < a < b < +\infty$. Now consider a rectangle with vertices (in counterclockwise direction) a , b , $b + 2\pi i$ and $a + 2\pi i$. The image of this rectangle, mapped by e^z , is a concentric ring centered at the origin with radii e^a and e^b respectively, [5].

IV. PROBLEMS IN THE EMPIRICAL DATA COLLECTION

The sound sources like thunder, the military weapons (for example howitzer) and the bells are complicated sound sources with a very wide frequency range and a unique dynamic range of the transmitted signal. Their spectrum contains infrasound, sound and ultrasound partials. The dynamic range is very large too and it cannot be detected entirely by human ear whose dynamic range of perception is about 120dB.

The accuracy of the measurements is influenced by the meteorological factors:

- The air temperature and humidity, wind speed and direction, and pressure in slight degree;
- The turbulence characteristic: structural constants of temperature, wind velocity and acoustic refractive index fluctuations;
- The vertical distributions of the meteorological parameters and characteristics.

The ambient noise and underlying surface (its characteristics determining the sound absorption by the ground and sound reflection and the relief) are very significant factors too.

The best of all over the world measurement set with a corresponding measurement microphone at this moment was used because of this [6]. For example, this set is able to process the signal without distortions with dynamic range up to 160dB.

Experimental set-up includes:

- Pressure-field Microphone Type 4193 Brüel&Kjær, [6] available in Transducer Electronic Data Sheet (TEDS) combinations with the classical Preamplifier Type 2669 with an individual calibration; Dynamic Range: 19 ... 162 dB, Sensitivity: 12.5mV/Pa.
- Vibration Transducer TRV-01 SPM Instrument;

- Compact Data Acquisition Unit 3560B Brüel & Kjær, [6] for outdoor use that consist: Dyn-X input modules with a analysis range exceeding 160 dB and automatic detection of front-end hardware and transducers – supports IEEE 1451.4-capable transducers with TEDS (Transducer Electronic Data Sheet); output TCP/IP protocol communication - RJ 45 connector complying with IEEE-802.3100baseX; Multiframe Control option;
- Base software PULSE 12 for CPB (Constant Percentage Band) analysis 2 channels; 5-channel Time Capture; PULSE Bridge to MATLAB®
- MathWorks Software - MatLab&Simulink, toolboxes for FFT and Wavelet analysis.

It can be seen, that the hardware equipment and the software are very qualitative. A part of equipment is shown in Fig.2.

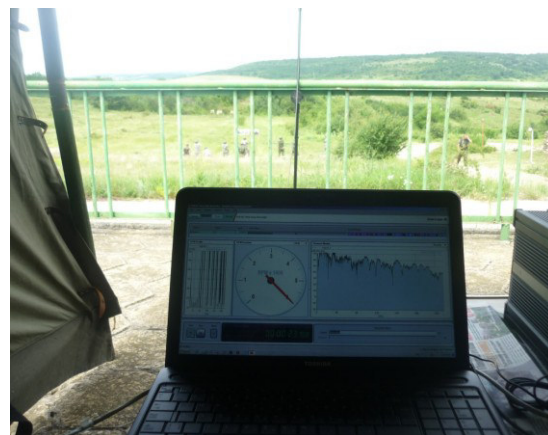


Figure 2. Part of measurement equipment

V. ANALYSIS AND DISCUSSION OF RESULTS

On the figures 3-11 the signals of thunder, volley of 122mm self-propelled howitzers and bell are shown, as well as and their soundprints. It can be seen that all prints (respectively scalograms) provide very rich information. Conclusion is that the signals have a complicated structure. In case of thunder it can be due from fractal nature of the lightening. It is known that the gun shot produces two main wave fronts: the shock wave and the muzzle blast. They can be located on the respective figures.

The soundprint of the bell (untuned orthodox bell) possess similar features. It is clear because of the fact that the bell is 3D impulse source.

But the sound multiplicity (complexity) in the bells pitch sensing is present all the time and there is some ambiguity of its tonality. Notating the bell music by ear is linked with a range of acoustic hindrances, unlike the music of other musical instruments, [7].

On the figures 3-5 are illustrated the thunder's waveform in time, CWT and soundprint. This signal is collected by microphone 4193 Brüel & Kjær and PULSE data

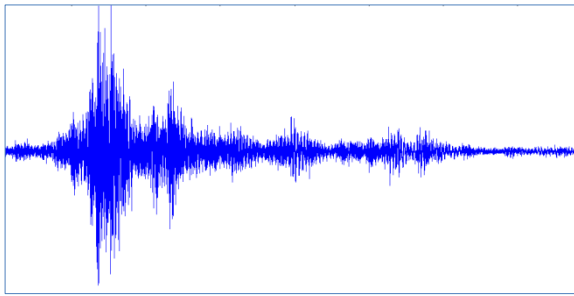


Figure 3. Record of a thunder

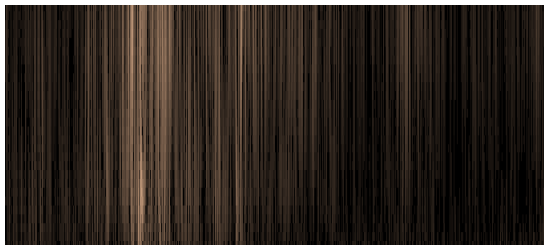


Figure 4. Scalogram of a thunder



Figure 5. Soundprint of a thunder

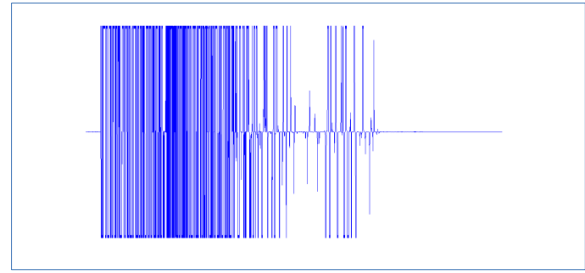


Figure 6. Record of a volley of two 122mm howitzers

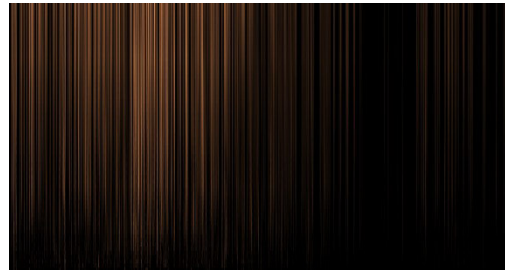


Figure 7. Scalogram of a howitzer's volley

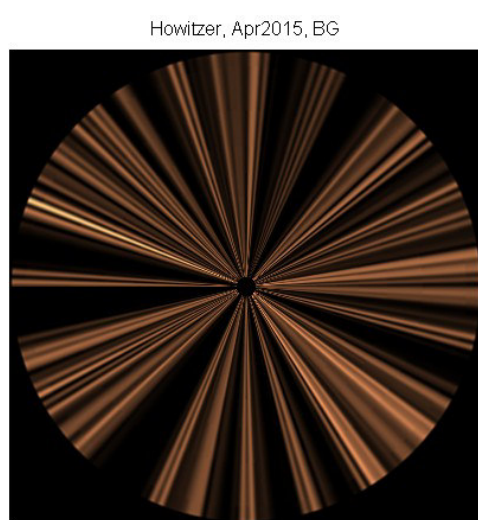


Figure 8. Soundprint of a howitzer's volley

acquisition unit 3560B, with sample frequency $F_s = 2^{16}$ samples per sec. (the microphone has wide dynamic range, no distortions).

On the figures 6-8 are illustrated the volley waveform of two 122-mm self-propelled howitzers (*portions of the signal are cut off* - signal is recorded by the notebook's internal microphone).

In the figures 9-11 was shown the bell strike waveform in time and its scalogram and conformal map (soundprint). This signal is collected by microphone 4193 Brüel & Kjær and PULSE data acquisition unit 3560B, with sample frequency $F_s = 2^{16}$ samples per sec.

The scalogram and the soundprints are applicable to estimation of the beat frequencies and modal damping ratios of unique bells,[8,10].

The modern web based approach to managing audio and video archive gives the opportunity to work with multimedia database, [9].

VI. CONCLUSIONS AND FUTURE WORKS

The results of researches obtain, that at some natural phenomena and manmade sounds there is a high level of infrasound component.

For obtaining system of attributes of specific noises, the authors use methods of conformal transformation. This approach allows receiving the optimal structure of the signals.

In the future the authors plan to create a large data base for detecting, identification and localization of difficult to research acoustic phenomena. The results will be useful in civil and military domains.

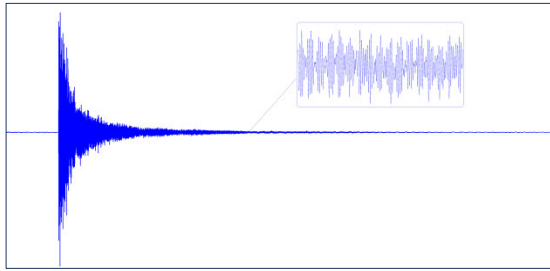


Figure 9. Record of the Bulgarian bell stroke (XIII cent.), [10]



Figure 10. Scalogram of the bell stroke

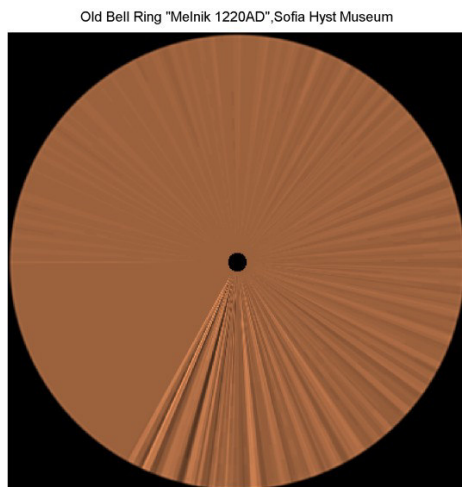


Figure 11. Soundprint of the bell

REFERENCES

- [1] Alissa Greenberg, Google Lost Data After Lightning Hit Its Data Center in Belgium, Aug. 20, 2015, online at „<http://time.com/4004192/google-data-lightning-belgium/>”
- [2] Dimitra Karamanides, *Pythagoras: pioneering mathematician and musical theorist of Ancient Greece*, The Rosen Publishing Group, 2005
- [3] Daubechies, Ingrid, *Ten lectures on wavelets* (CBMS-NSF regional conference series in applied mathematics; 61), SIAM 1992
- [4] Lars Ahlfors, *An Introduction to the Theory of Analytic Functions of One Complex Variable*, New York, 1953.
- [5] Dimkov G., T. Trifonov, I. Simeonov „Improving the visualization of scalograms by means of transformation in the complex plane”, *Proceedings of the Forty Second Spring Conference, of the Union of Bulgarian Mathematicians*, Borovetz, April 2–6, 2013, p.218. „http://www.math.bas.bg/smb/2013_PK/tom_2013/pdf/218-221.pdf”
- [6] Brüel&Kjær, www.bksv.com
- [7] Aldoshina, I.A. and Nicanorov, A "The Investigation of Acoustical Characteristic of Russian Bells", *Audio Engineering Society Convention 108*, February 19-22, 2000, Paris, France
- [8] Trifonov T., I. Simeonov, R. Dzhakov „Features of time-frequency analysis visualization of large dynamic range signals”, *XLVII International Scientific Conference on Information, Communication and Energy Systems and Technologies (ICEST 2012)*, Veliko Tarnovo, June 2012, pp. 155-158
- [9] Tsvetanka Georgieva-Trifonova, *Databases*, Astarta, 2010 (in Bulgarian)
- [10] The Bell Project-Research and Identification of Valuable Bells of the Historic and Culture Heritage of Bulgaria and Development of Audio and Video Archive with Advanced Technologies, <http://www.math.bas.bg/bells/belleng.html>

Measurements with a Non-Invasive Pulse oximetry Module

Bertalan Beszédes*, Tamás Egressy**

* Óbuda University/ Alba Regia Technical Faculty, Székesfehérvár, Hungary

**Albacomp RI Kft, Székesfehérvár, Hungary

beszedes.bertalan@amk.uni-obuda.hu, tamasegressy@gmail.com

Abstract— The pulse oximeters are commonly used devices in medical practice. With their help, the oxygen saturation of the haemoglobins, or in other words, the oxygen saturation of arterial blood can be measured with a few percentage of error. The pulse oximeters can measure continuously, in a non-invasive way. By monitoring the oxygen saturation, reduced oxygen supply can be diagnosed. The medical term of this condition is hypoxemia. The most common use of pulse oximeters is to help doctors diagnose hypoxemia.

I. INTRODUCTION

Normal blood oxygen levels in human arteries are considered between 95-99%. This means that the average percentage of haemoglobin molecules that deliver oxygen is 97%. In other hand, this percentage in the veins is only 75%. It's irrelevant, if someone runs, or sleeps, only the blood flow of blood vessels change, the oxygen saturation remains constant. Reduced oxygen saturation is usually a result of an abnormal physiological process.

In medical practice, pulse oximeters are used to monitor the patient in numerous cases: during anaesthesia, child-birth, emergency intervention, operation, or in case of cardiovascular problems, or pulmonary disease, but pulse oximeters are in everyday use in hospital rooms, during patient transport, and in maternity ward.

Pulse oximeter sensors, hereinafter SPO2 sensors can be found in everyday life in the form of wearable bracelets. They can be a really useful tool in forecasting diseases or emergencies.

II. THEORY

A. Hemoglobin

Haemoglobin is a type of protein that is carried by red blood cells. The proportion of blood that is occupied by red blood cells is usually 45%. The haemoglobin's main role is to deliver oxygen from the lung to the tissues, and to carry the end products of metabolism, such as carbon-dioxide (CO₂) from the tissues to the lung. The haemoglobin protein contains iron molecules, that can bind oxygen molecules. This process is called oxygenation. This term is not the same as oxidation. One iron molecule is able to bind four oxygen molecules. In this case, the oxygen saturation is 100%. A fully saturated haemoglobin is called oxyhaemoglobin, its short form is HbO₂. The fully depleted, oxygen-free haemoglobin is called deoxyhaemoglobin, in short form: Hb.

Haemoglobin gives the distinctive red colour of the blood. The oxygen-rich blood in arteries is bright red, the oxygen-poor blood in veins has a colour between brown

and dark red, observed through the skin, it's blueish. It can be observed that the haemoglobin's colour, or in other words the absorption spectra is closely related to its oxygen saturation. At a given wavelength, the haemoglobin, depending on its oxygen saturation absorbs more or less light. This property of the haemoglobin is used to measure blood oxygen levels. The oxyhaemoglobin and deoxyhaemoglobin's absorption spectra can be seen on fig.1. On the left side the wavelength is 450-1000 nm, on the right side 650-1000 nm is highlighted.

The human eye is sensitive to light in the interval of

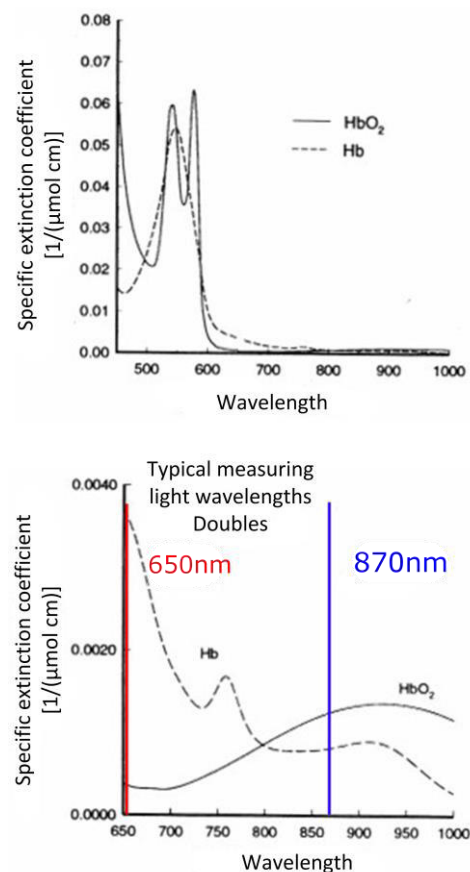


Figure 1. The absorption spectrum of hemoglobin [1]

light in the interval of 400-700 nm. The oxyhaemoglobin and deoxyhaemoglobin's absorption curves can be distinguished on different wavelengths. It can be seen that the difference between the curves is greater by an order of magnitude in wavelengths below 600 nm. On the first

glance, it would seem like a better alternative to measure absorption, however it's an incorrect assumption.

B. The „Oximetry Window”, or the decisive factors of the measured light's wavelength

The 650-1000 nm interval in the electromagnetic spectrum of the light is the most suitable for oxymetric measurements. It's causes will be discussed in the following.

The upper limit: The water's absorption:

An adult human's body consists of 60% of water. The water is present in every tissue, such as in blood. Around 900nm the absorption of water starts to rise as the wavelength increases. When the wavelength reaches 1400 nm, the absorption is greater by an order of magnitude compared to absorption of the 650-1000 nm interval. The absorption of haemoglobin is relatively small, when measurement takes place, the light's wavelength can't go above 900-1000 nm, because the water would absorb a large part of the light instead, making the signal-to-noise ratio unacceptable. [7]

The lower limit: The skin's absorption:

In the 450-600 nm interval, the oxyhaemoglobin and deoxyhaemoglobin's absorption are well separated (fig.1 left side), this way, changes in oxygen saturation are well distinguishable. This would result in a better signal-to-noise ratio than the 650-1000 nm interval, but again, it's an incorrect assumption. In 450-600 nm interval, the measurement is prevented by the skin itself. The human skin has a great absorption of light in 400-650 nm interval, but above that, the skin's absorption decreases greatly, therefore using light above 650 nm wavelength for measurements is recommended. [5], [6]

C. Determining oxygen saturation

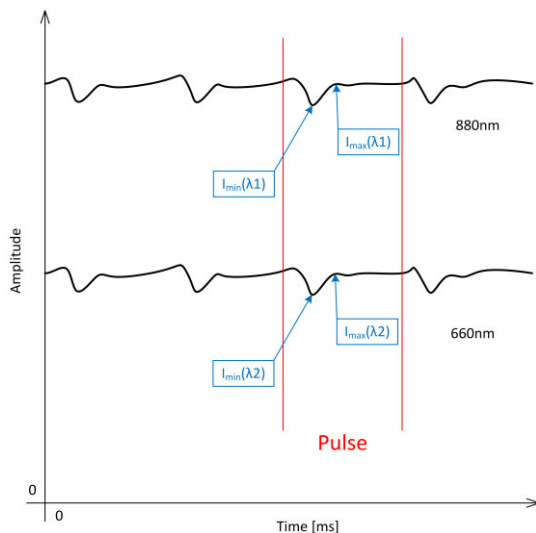


Figure 2. Typical pulse oximeter probe waveform used in two wavelengths

In pulse-oximeter devices, at least two, different wavelength light sources must be present. On Fig.2 two measuring light wavelengths in the absorption spectra of haemoglobin are displayed.

Two specific measuring light wavelengths can be seen on Fig.1 in the absorption spectrum of haemoglobin. In case of 870 nm, with the decrease of oxygen saturation, the curve moves towards the dashed line, that is, the absorption decreases. This way the light's intensity, that passes through the blood vessels, decreases. In case of 650 nm things work the opposite way, since the oxyhaemoglobin's absorption is lower than the deoxyhaemoglobin's.

The phrase pulse oximeter is originated from the pulsating nature of arterial blood. Every tissue has constant light absorption properties, only the arteries pulsate. Each heartbeat means that the fresh, oxygen-rich blood replaces the oxygen-poor blood, therefore changes the artery's light absorption. If tissue interwoven with arteries is illuminated with two different wavelength light sources, and the light is detected and converted into electrical signs, we can get a figure similar to Fig.2.

The arterial pulsation modulates the light passing trough. The large DC component is a result of the absorption of tissues, therefore they are not pulsating. I_{max} and I_{min} are in compliance with the systolic and diastolic states. The effect of the non arterial components can be eliminated by separating and amplifying the pulsating component, therein lies pulse-oxymetry's greatest advantage.

The dominant light absorbent material in 650-1000 nm range is the oxyhaemoglobin and deoxyhaemoglobin, therefore it is sufficient to write down the equation of the oxygen saturation with only these components.

The Lambert-Beer Law takes the form of the following Equation (1): [2], [3], [8]

$$I = I_0 e^{-\sum \alpha_i c_i l_i} \quad (1)$$

Where I₀ is the intensity of the light entering the tissue, that is emitted by the LEDs. The measurement of this intensity would be too complicated, so we look on its magnitude as an unknown. a_i and c_i are absorption coefficients and concentrations of tissue-components, l_i denotes the distance traveled by light in said tissue-component. Measuring l_i is also problematic, therefore let's consider it as an unknown.

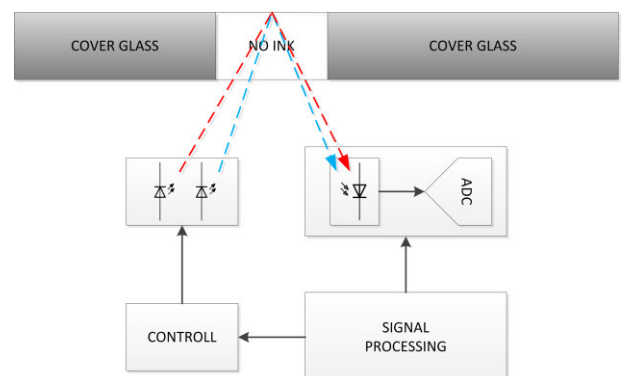


Figure 4. Introducing the principle of Measurement

Next, we have to write down Eq. (2) on one pulse's systolic and Eq. (3) diastolic points (see on fig. 2) Assuming the artery's thickness, and thereby the distance

travelled by the light in the artery changes by λl between the systolic and diastolic points. On this route, the attenuation can only be caused by the arterial blood, therefore we can write down the HbO_2 and Hb 's absorption coefficients and concentrations next to λl . $I_{\max}$ and $I_{\min}$ correspond to figure above.

$$I_{\max} = I_0 e^{-\sum \alpha_i c_i l_i} \quad (2)$$

$$I_{\min} = I_0 e^{-\sum \alpha_i c_i l_i + \Delta l (\alpha_{Hb} c_{Hb} + \alpha_{HbO_2} c_{HbO_2})} \quad (3)$$

Variables indexed by "i" describe the impact of tissues that are irrelevant from the perspective of pulse oxymetry, moreover they include attenuation of the blood found in a relaxed arteries. Two equations quotient, and logarithm takes the form of the following Equation (4):

$$\ln \frac{I_{\max}}{I_{\min}} = \Delta l (\alpha_{Hb} c_{Hb} + \alpha_{HbO_2} c_{HbO_2}) \quad (4)$$

It is clear to see that the entry intensity, generated by the LEDs is not needed, this way the first unknown is eliminated. The distance travelled by the light (Δl) remains an unknown. It can be eliminated by measuring on two, different wavelengths. This is why two different light sources are needed. The absorption coefficient depends on the wavelength, and accordingly the indexes λ_1 and λ_2 indexes are displayed accordingly.

The two wavelength equation's quotient is the so called "R-ratio" (Ratio of the ratios) Eq. (5). In this equation, the distance travelled by the light is eliminated. This ratio is measured by the pulse-oxymeters. The R-ratio is the ratio of the amplitude of two pulsating components, with different wavelengths. The arterial oxygen saturation can be determined from the R-ratio. [3]

$$R = \frac{\ln \frac{I_{\max}(\lambda_1)}{I_{\min}(\lambda_1)}}{\ln \frac{I_{\max}(\lambda_2)}{I_{\min}(\lambda_2)}} = \frac{\Delta l (\alpha_{Hb\lambda_1} c_{Hb} + \alpha_{HbO_2\lambda_1} c_{HbO_2})}{\Delta l (\alpha_{Hb\lambda_2} c_{Hb} + \alpha_{HbO_2\lambda_2} c_{HbO_2})} \quad (5)$$

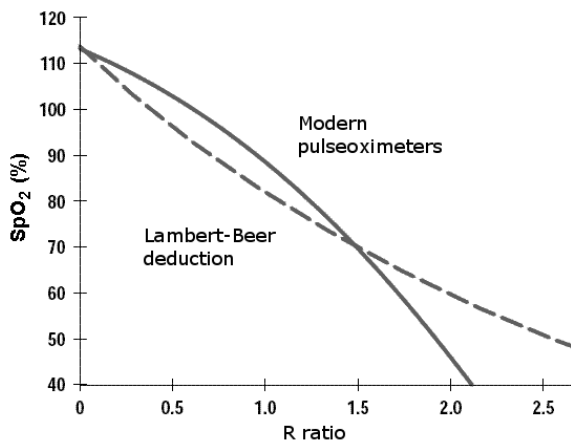


Figure 3. Modern pulse oximetry measurement graph [4]

Where the SpO_2 is the saturation of oxygen in the blood, c_{HbO_2} is the oxygenated hemoglobin concentration and c_{Hb} is the deoxygenated hemoglobin concentration.

The Lambert-Beer Law's boundary conditions are never totally met in practice. The model does not take into account the multi-scattering- phenomenon on red blood cells, the inhomogeneous tissue structure, and the fact that the light emitted by the LEDs is not monochrome. Therefore today the correlation between the R-ratio and SpO_2 -value is not determined with the help of the absorption coefficients. Instead it is deduced from hundreds of SpO_2 measurements of volunteers, with statistical methods. After blood sampling, the samples undergo a blood gas analysis, accurate SaO_2 concentration is determined, after that, the results are compared to the R-ratio measured by the pulse-oxymeter. The result of this process is called "Calibration-table". After the measurement of R-ratio, the oxymeter reads from the pre-programmed Calibration-table [7]

III. A INTRODUCTION OF HARDWARE

Msp430 is a perfect microcontroller for handling SpO_2 measurements, it even has resources to handle additional peripherals, such as OLED display, bluetooth module, thermometer. These can be the elements of a smartwatch.

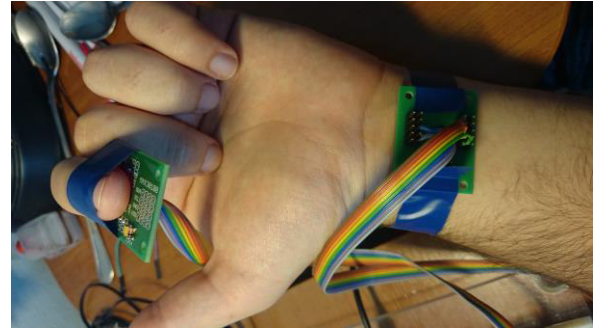


Figure 7. Measuring with the devices

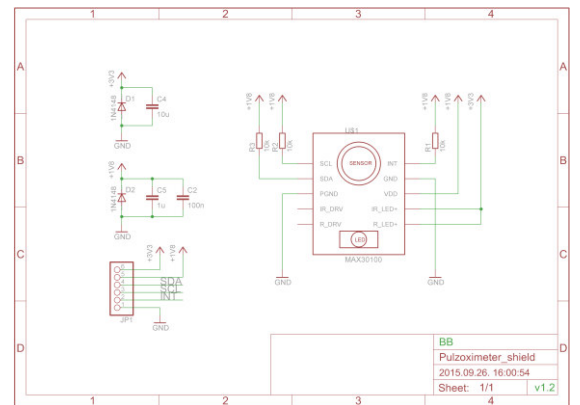


Figure 6. Schematic of the device

The operating principle of the sensor is shown on Fig. 4 and Fig. 5. The device combines an infra-red and a red LED light with a photodetector. It includes an optimized lens and low noise analog signal processing unit. The sensor is able to communicate with a microcontroller through an I2C bus (Fig.6).

The results shown in Fig. 8 and Fig. 9 were measured with the measurement setup shown in Fig.7. The measurements provide similar results on fingertips, and on wrists as well. The results mainly differ in amplitude, but heart rate periods are well distinguishable.

continuous monitoring of cardiac arrhythmias and other cardiac events.

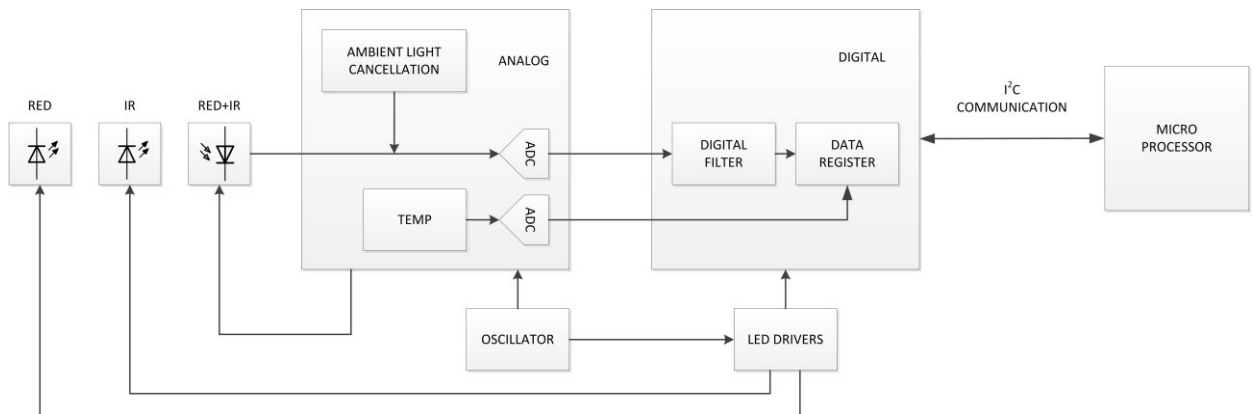


Figure 5. The device structure

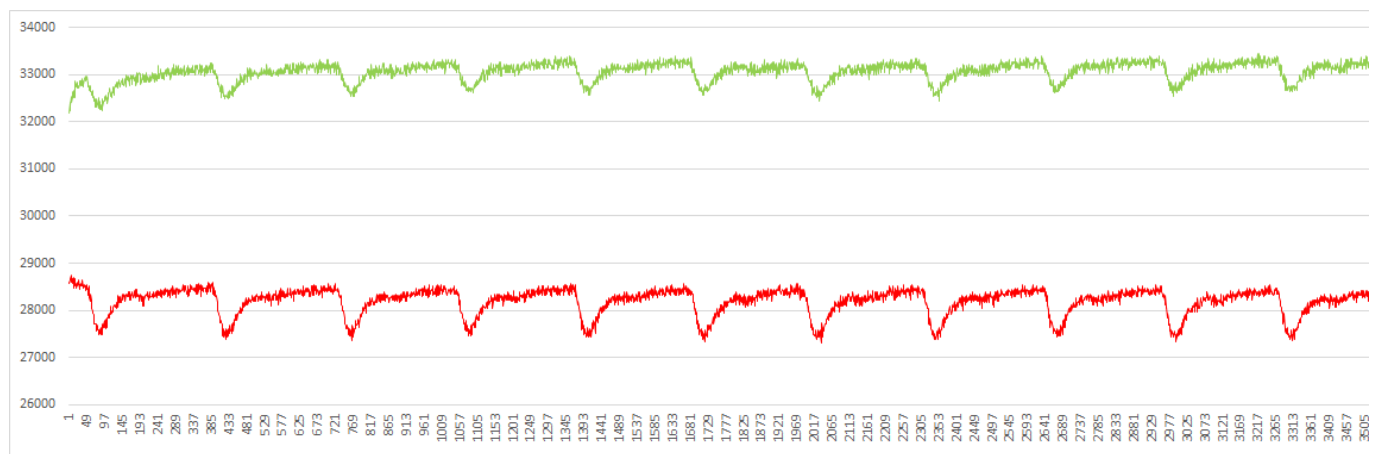


Figure 8. Measuring on fingertip

IV. CONCLUSION

Generally pulse oximeters are used in hospitals, for monitoring patients. Our measurement results showed that they are also able to produce good results on wrist, therefore, pulse oximeters can be introduced in everyday life, for instance for fitness-wellness purposes.

The actual heart rate can be calculated, using the results of the measurements. Measurements can be also used for

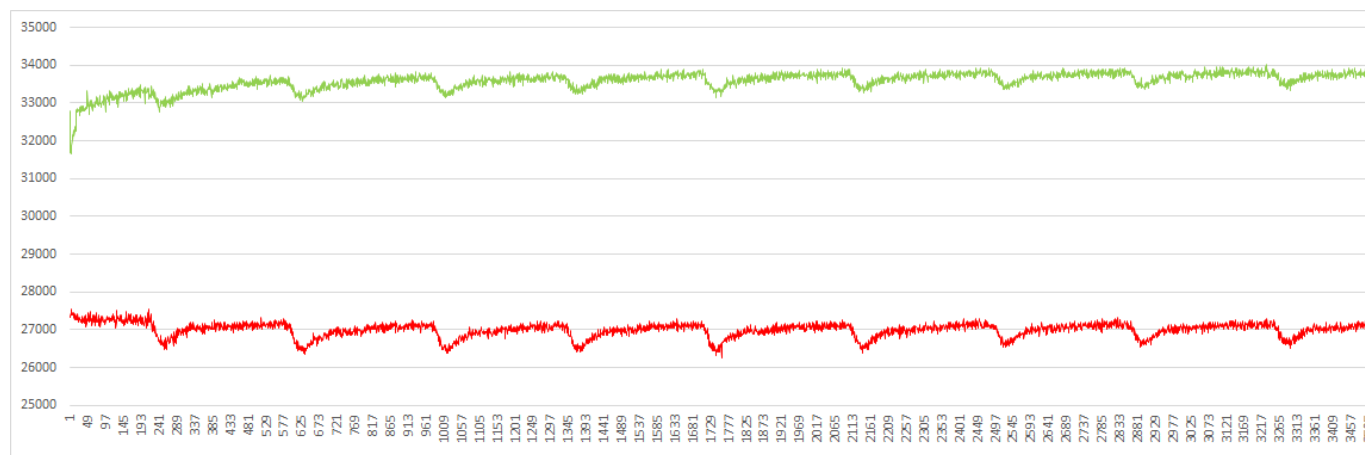


Figure 9. Measuring on wrist

REFERENCES

- [1] Horecker B. L.: The absorption spectra of hemoglobin and its derivatives in the visible and near-infrared regions, *J. Biol. Chem.*, 1943, 148: 173-183. J. Clerk Maxwell, *A Treatise on Electricity and Magnetism*, 3rd ed., vol. 2. Oxford: Clarendon, 1892, pp.68-73.
- [2] Bohren C.F., Huffman D.R.: Absorption and Scattering of Light by Small Particles, Wiley-Interscience, New York, 1983; Cope 1991 K. Elissa, "Title of paper if known," unpublished.
- [3] Kästle S., Noller F., Falk S., Bukta A., Mayer E., Miller D.: A New Family of Sensors for Pulse Oximetry, *Hewlett-Packard Journal*, 1997 February, Article 7., pp. 3-4. Y. Yorozu, M. Hirano, K. Oka, and Y. Tagawa, "Electron spectroscopy studies on magneto-optical media and plastic substrate interface," *IEEE Transl. J. Magn. Japan*, vol. 2, pp. 740-741, August 1987 [Digests 9th Annual Conf. Magnetics Japan, p. 301, 1982].
- [4] Kästle S., Noller F., Falk S., Bukta A., Mayer E., Miller D.: Volunteer Study for Sensor Calibration, *Hewlett-Packard Journal*, 1997 February, Subarticle 7a, pp. 1-3. Marbach 1993
- [5] Winey B., Yu Y.: Oximetry considerations in the small source detector separation limit, 2006, Conf. Proc. IEEE Eng. Med. Biol. Soc., 1:1941-3.
- [6] Friedland S., Benaron D., Parachikov I., Soetikno R.: Measurement of mucosal capillary hemoglobin oxygen saturation in the colon by reflectance spectrophotometry, *Gastrointestinal Endoscopy*, 2003, Vol. 57, No. 4, pp. 492-497.
- [7] Marbach R., Koschinsky T., Gries F. A., Heise H. M.: Noninvasive Blood Glucose Assay by Near-Infrared Diffuse Reflectance Spectroscopy of the Human Inner Lip, *Appl. Spectrosc.*, 1993, 47, pp. 875-881.
- [8] Stubán Norbert, Vezeték nélküli magzati pulzoximéter megvalósítása, Budapest, 2009, pp.18-23.
- [9] Cope M.: The application of near infrared spectroscopy to non invasive monitoring of cerebral oxygenation in the newborn infant, PhD thesis, Dept. of Medical Physics and Bioeng., University College London, 1991.
- [10] Gy. Györök - Function monitoring of humane body by 3d a-sensor - BULETINUL STINTIFIC AL UNIVERSITATII POLITEHNICA DIN TIMISOARA ROMANIA SERIA AUTOMATICA SI CALCULATORAE / SCIENTIFIC BULLETIN OF POLITEHNICA UNIVERSITY OF TIMISOARA TRANSACTIONS ON AUTOMATIC CONTROL AND COMPUTER SCIENCE 51(65):(1) pp. 61-64. (2006)
- [11] Gy. Györök, M. Makó, J. Lakner - Combinatorics at Electronic Circuit Realization in FPAA - 2008 6TH INTERNATIONAL SYMPOSIUM ON INTELLIGENT SYSTEMS AND INFORMATICS. - Konferencia helye, ideje: Szabadka, Szerbia, [2008.09.26-2008.09.27](#). New York: IEEE, 2008. pp. 1-4. (International Symposium on Intelligent Systems and Informatics) (ISBN: [978-1-4244-2406-1](#))
- [12] Gy. Györök - Self Organizing Analogue Circuit by Monte Carlo Method - Proceedings of the International Symposium on Logistics and Industrial Informatics (LINDI 2007). Konferencia helye, ideje: Wildau, Németország, [2007.09.13-2007.09.15](#). Budapest: IEEE Hungary Section, 2007. pp. 1-4. (ISBN: [1-4244-1441-5](#))

Land Cover Change Detection by Remote Sensing

M. Verőné Wojtaszek*

* Institute of Geoinformatics, Alba Regia Faculty of Engineering, Óbuda University, Székesfehérvár, Hungary
wojtaszek.malgorzata@amk.uni-obuda.hu

Abstract— Land cover mapping is one of the most important and typical applications of remote sensing data. Land cover corresponds to the physical condition of the ground surface, for example, forest, grassland, concrete pavement etc., while land use reflects human activities such as the use of the land, for example, industrial zones, residential zones, agricultural fields etc. Land cover change detection is necessary for updating land cover maps and the management of natural resources. The change is usually detected by comparison between two multi-date images, or sometimes between an old map and an updated remote sensing image. Information on land cover and changing land cover patterns is directly useful for determining and implementing environment policy and can be used with other data to make complex assessments (e.g. mapping erosion risks). In this paper we focus on change detection techniques provided by IDRISI and eCognition software. Optical and multispectral methods include pixel-based and object-oriented approaches, and approaches that deal with changes between optical images and existing vector data. What method to use is dependent on the sensor, the size of the changes in comparison with the resolution, their shape, textural properties, spectral properties, and behavior in time, and the type of application.

I. INTRODUCTION

Land cover changes continuously. The rate of change can either be dramatic and unexpectedly sudden, such as the changes caused by hurricanes, fire, or subtle and gradual, such as regeneration of forest, soils. Numbers of studies have described that land cover has changed dramatically during the past several centuries and that these changes have affected our ecosystems [1]. The causes of land cover changes are divided into five groups [2]:

- long-term natural changes in climate conditions
- geomorphological and ecological processes
- human-induced alterations of vegetation cover and landscapes
- interannual climate variability
- human-induced greenhouse effect

The analysis of land use and land cover change plays a key role in modern landscape research and landscape ecology. Different fields such as geographical information science and remote sensing have contributed considerably to this area [3], [4]. The change is usually detected by comparison between two multi-date images, so the length of a considered analysis period is limited to the time span for which remotely sensed data (aerial photography,

satellite images) is available. To consider longer time periods, historical maps can be used [5]. Comparing these maps with contemporary maps requires an extensive evaluation since they usually contain considerable inherent uncertainty.

Land cover change detection is necessary for updating land cover maps, accurate and up-to-date information is needed for many applications. One of the objectives of land cover change detection is to understand relationships and interactions between humans and the environment in order to manage and use resources in a better way for sustainable development [4]. Information on land cover and changing land cover patterns is directly useful for determining and implementing environment policy and can be used with other data to make complex assessments (e.g. mapping erosion risks).

The number of remotely sensed dataset available for analysis has increased markedly since 1972, when LANDSAT data became available. The archive of LANDSAT datasets are especially useful for monitoring long-term ecosystem effects. Tools and techniques are needed to detect, describe, and predict these changes. During the last decades many change detection techniques have been developed and applied to assess land cover changes. The goal of this papers is to summarize some important aspects of change detection, including data selection, and detection tools and methods within IDRISI and eCognition software.

II. DATA SELECTION

The success of applying remotely sensed data for change detection depends on careful selection of the data source. The important attributes of images can be taken into consideration are spatial, temporal, spectral and radiometric resolution.

III. LAND COVER CHANGE DETECTION

There are two basic types of remote sensing change detection. The first one is map-to-map and the other is image-to-image comparison. In map-to-map comparison, individual land cover maps are created independently using different data of images, and the results are compared. In this paper are showed change detection techniques provided by IDRISI and eCognition software.

A. Land Change Modeler

The chapter is dealing with the results of the author long term researches on land use changes of Lake Velence watershed.

The Land Change Modeler is a vertical application within IDRISI oriented to the pressing problem of accelerated land conversion and the very specific needs of biodiversity conservation. The LCM application provides a robust set of tools for the analysis of change and the creation of viable plans and scenarios for the future. The major tasks of the Land Change Modeler interface are: Change Analysis (analyzing past land cover change), Transition Potentials (modeling the potential for land transitions) and Change Prediction (predicting the course of change into the future) [6].

Objectives of the study were the following: land use change in the last two decades (using remote sensing to classify land use), land change analysis (which land use changes have recently occurred and where), determine the landscape characteristics influencing the spatial pattern of the land use transitions, and to predict future land use change using information from past changes patterns.

LANDSAT TM (1990, 2006), SPOT (2000) and ASTER (2004) were interpreted and classified dynamic change analysis (spatial pattern of the changes; resolution 30m, 15m, 10m).

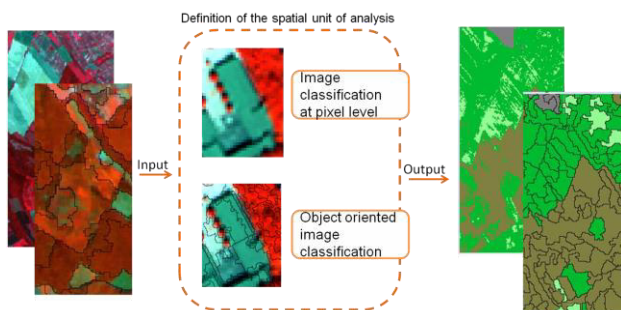


Figure 1. Land cover information extraction from satellite images

The land cover maps were created using supervised classification techniques (pixel- and segment-based) in IDRISI TAIGA. The classification was performed on multi-channel data and this process assigns to each pixel or segment of an image to a particular class based on statistical characteristics of the pixel (or segment) values (Fig.1) [7].

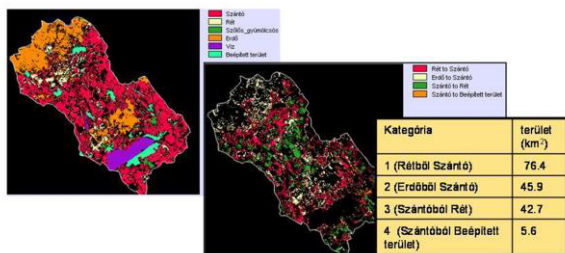


Figure 2. Land change analysis using LCM

The LCM provides a set of tools for analyzing land use (cover) change, including graphs of gains and losses, net changes and contributions experienced by any category. A simple one-click interface provides the ability to generate rapid maps of change, persistence, specific transitions and

exchanges between categories. The analysis of two land use maps of different dates showed that between 1990 and 2004 the amount of urban land, rural cover types of agriculture and water class increased (Fig.2). The urban area increased by 13% of the total area, the agriculture by 15% and the water by 8%. The water area increase is dependent on drought. The meadow field, the forest category and vine-orchard decreased. Decreasing of the forest is connected with forestry and means land cover change, but not land use.

B. Spectral Indices

Spectral indices derived from satellite data are widely used for land cover change research. They can reduce the data volume for analysis and provide combined information that is more strongly related to changes than any single band [4].

Vegetation indices use various combinations of multi-spectral satellite data to produce single images representing the amount of vegetation present, or vegetation vigor. Low index values usually indicate little healthy vegetation while high values indicate much healthy vegetation. Different indices have been developed to better model the actual amount of vegetation on the ground.

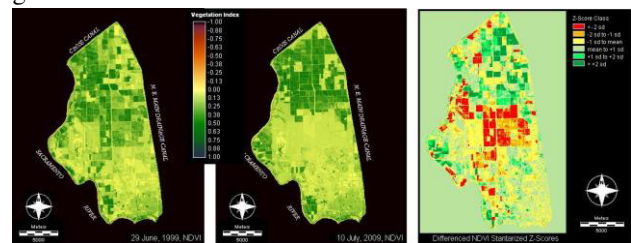


Figure 3. Land change analysis using vegetation indices [8]

There are many quantitative methods we can use to analyze change between images. One of them is simple differencing, when we merely subtract one images from the other, then analyze the result. Figure 3 illustrates the appearance of NDVI images and differencing image. The values are based on the multi-temporal NDVI values. The result values (Z-Scores = ((observed value - mean value)/standard deviation)) are calculated from pixel values, mean, and standard deviation values. The values represented with negative standard deviations (std) represent negative vegetation changes between 1999 and 2009. Positive values represent positive changes [8].

C. Change Detection, working with "maps" in eCognition

A successful change detection system has to analyze the change of each phenomenon in its characteristic scale and then fuses the resulting change maps into a single change map. Many studies made use of multi-scale analysis in literature. Multi-scale object specific analysis is a multi-scale approach that automatically defines unique spatial measures specific to individual image-objects composing a scene. These specific measures are used in a weighting function to automatically upscale an image to coarser resolution [9, 10].

Recently, a new system for multi-scale data set generation is implemented in eCognition software. This approach extracts objects of interest at the scale of

interest, segmenting images by operating on the relationships between linked objects. The multi-scale change detection system consists of the following phases; superimposing the two-date images; multi-scale data-set generation, individual scale change detection, optimal scale detection, and scale-based change fusion.

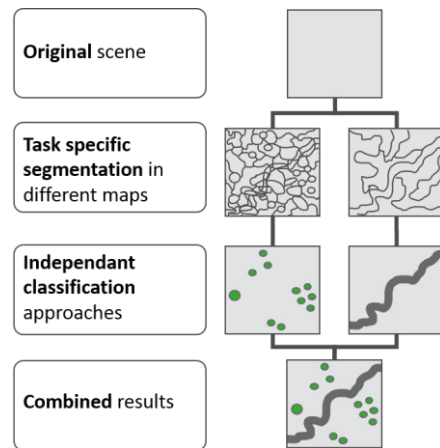


Figure 4. Change detection, working with „maps” in eCognition [11].

IV. CONCLUSION

In this paper we focus on change detection techniques provided by IDRISI and eCognition software. Optical and multispectral methods include pixel-based and object-oriented approaches. What method to use is dependent on the sensor, the size of the changes in comparison with the resolution, their shape, textural properties, spectral properties, and behavior in time, and the type of application.

The application of the multi-scale detection system ensures less false change alarms. The study of individual objects is performed in the object's preferred scale so that noise is rejected. The detected change objects keep their true shapes without any distortion on their boundaries or

inside the objects that may result from noise or image sampling. Applying multi-scale change detection provides an advantage in the overall accuracy over the single scale change detection.

REFERENCES

- [1] G. M. Foody “Assessing the accuracy of land cover change with imperfect ground reference data,” *Remote Sensing of Environment*, vol. 114, pp. 2271–2285, 2010.
- [2] E. F. Lambin and A. H. Strahlers, 1994): “Indicators of land cover change for change-vector analysis in multi-temporal space at coarse scales” *International Journal of Remote Sensing*, vol. 15. pp. 2099-2119, 1994
- [3] D. Lu et al., “Change detection techniques”, *International Journal of Remote Sensing*, vol. 25. pp. 2365-2401, 2004
- [4] P. Coppin et al. “Digital change detection methods in ecosystem monitoring” *International Journal of Remote Sensing*, vol. 25. pp. 1565-1596, 2004
- [5] M.. Verone Wojtaszek and L. Roncyk, Object-based Classification of Urban Land Cover Extraction Using High Spatial Resolution Imagery, *International Scientific Conference on Sustainable Development & Ecological Footprint (NymE TÁMOP 4.2.1/B)*, Proceedings, ISBN 978-963-334-047-9, 7 pp. 2012
- [6] IDRISI Taiga. GIS/Image processing software, World Wide Web homepage <<http://www.clarklabs.org/>>
- [7] M. Veróné Wojtaszek “Földhasznosítás változásának követése távérzékeléssel a Velencei-tó vízgyűjtőjében” *Geomatikai Közlemények*: pp. 273-280. 2009
- [8] J. Alvares “Land Cover Change Analysis” 2009, <http://academic.emporia.edu/aberjame/student/alvarez4/natchange.html>
- [9] G. Hay, D. Marceau, Multiscale Object-Specific Analysis (MOSA): An Integrative Approach for Multiscale Landscape Analysis, *Remote Sensing and Digital Image Processing (4): Chapter 3*. 2004
- [10] G. Hay,T. Blaschke, D. Mareceau, A. Bouchard, A Comparison Of Three Image-Object Methods For The Multi-Scale Analysis Of Landscape Structure , *ISPRS Journal for photogrammetry and remote sensing (57)*: 327-345, 2003.
- [11] eCognition tutorial: New OBIA dimensions: Maps

The Impact of Alba Regia Information Technology Competition for Alba Regia Technical Faculty's Enrollment

Dezső Hatalyák and Éva Hajnal

Óbuda University, Alba Regia Technical Faculty
Budai Str 45, Székesfehérvár, Hungary (H-8000)
hatalyak.dezso@amk.uni-obuda.hu
hajnal.eva@amk.uni-obuda.hu

ABSTRACT – An applied information technology competition was organized fifth time last year by the Alba Regia Technical Faculty's teacher and students. This competition is organized for regional students, who are potentially continuing their studies in our faculty. We compared the educational results of two groups with ten-ten students. The first group members are our university students, who were participants in the final few years earlier, and the second group members are randomly chosen from university students' sample. Our question is if the education of applied information technology is capable for developing intellectual abilities, that are required for further studies, and if the applied information knowledge is correlating with informatics thinking and capability, what we are expecting from our students. We compared the educational results of the two groups of students from different aspects. We used statistical hypothesis testing. It was found, that the progress of study, the results of programming course, and the grade point averages were significantly better in the group of competition participant students.

I. INTRODUCTION

An applied information technology competition for secondary school students was organized by Óbuda University's Alba Regia Technical Faculty for the fifth time [1][2]. For this jubilee anniversary raised an idea that we may reached enough information during this period for investigate the question, if this competition helped our faculty in the enrollment of students or not.

In the history of the Hungarian information technology education dominated two very different trends. In the first few years the information technology education was equal to teach programming. This fact was not accepted neither the parents nor the students, and this curriculum was not representative to the assembled IT knowledge. After that came a reform that had a wider social acceptance, but it made teaching windows applications into the center of information technology education. It get much criticism, that in the secondary school the teachers are forming secretaries, and this way of the information technology education is not capable to deliver the mentality and the ability development. We were obliged to organize our competition with respecting of the fact, that the information technology education in the secondary schools means applications only[2]. This fact posed some questions for us. First question is that the applier information technology skill is enough for developing intellectual ability? The second is if the competition mending our enrollment or not? It can help us to answer both questions, when we analyze the educational results of the students, who took part in this competition earlier, and after that chose our university. In our research the results of these students were examined, statistical analyses were executed, and explanations were searched onto the reasons of the differences.

II. DATA BASE AND CALCULATIONS

For the inspection we chose ten students, who took part in our competition earlier, when they were students in a secondary school, and after that they chose our university to their further education. Another ten students were chosen from our other students; they are taking part in engineering information technology BSc course or in information technology engineer course. We pick the members of the control group absolutely randomly. We analyzed the study outcome of the two groups each with ten persons. The data-base of the inspection was queried from the Neptun United Education System and from the competition results data base.

Table 1: Results of the students in our competition

Results of competitors			
Student	Result (points)	Maximum point	%
1 st	89	120	74,17%
2 nd	63	120	52,50%
3 rd	70	120	58,33%
4 th	51	110	46,36%
5 th	39	110	35,45%
6 th	51	110	46,36%
7 th	36	110	32,73%
8 th	48	120	40,00%
9 th	50	113	44,25%
10 th	45	120	37,50%

Table 2: Grades of the students who took part in the Alba Regia Information Technology competition

„md” means missing data that the student doesn't have that subject, because on the assistant of engineering of information technology training there is no Analysis I.

Student	Programming I.	Analysis I.	Average of all grades
1 st	4	4	3,84
2 nd	4	2	2,2632
3 rd	4	5	4,3
4 th	2	2	2,48
5 th	2	4	3,5167
6 th	5	5	4,8868
7 th	3	2	2,4118
8 th	3	3	3,2
9 th	3	4	3,7692
10 th	4	na	4,2

First the average grades in the two groups were calculated. For this we collated the marks of the students in both groups from the next subjects: Analysis I. and Programming I. You can ask, why these? In all training, which is related to information technology, the most difficult subjects are these, by our opinion. It is a real challenge when the students first time meet with programming, and mathematics at university level. Unfortunately, in the secondary school teachers teach generally neither the basics of programming. That's a very sad thing, because the students have to envisage that they have much deficiency in these areas of their knowledge. For the further calculations it was sup-

posed, that the marks of the students in our university have got a normal distribution.

Some idea for the interpretation of the results: „md” means that the student doesn't have that subject, because in information technology engineer course there is no Analysis I. By the technical manager BSc training there is no Programming I. We used instead Information technology lab subject, because these students first time in this subject meet with a programming language.

After we have collected data, first we calculated the averages, after that we used F-test for checking if the variances are equal or not. The T-test has to make in different way depending on if the variances are equal or not. With T-test we searched the answer, if there is a difference between the study result average, Programming I. and Analysis I. result of the two groups of students. Our null hypothesis was that there is no significant difference between the study result averages of the two groups[3].

Table 3: The grades of the students in the control group (randomly chose students)

Number of the student	Programming I.	Analysis I.	Average of all grades
1 st	2	3	2,6667
2 nd	4	5	3,746
3 rd	2	2	2,6667
4 th	2	2	2,6761
5 th	2	2	2,4643
6 th	2	2	2,2857
7 th	3	3	3,3846
8 th	3	3	2,2667
9 th	3	2	2,8387
10 th	2	2	1,5385

We investigated if there is a connection between the resulting points of the competition participant students and the posterior study-average, Programming I. and Analysis I. exam mark of these students. We made correlation calculation[3].

We collected how many times made a student a subject-repetition altogether and how many different subject had they to repeat. Table 4 shows this dataset.

Table 4: The subject-repeat statistic of the student who took part in our competition

Number of the student	Different subject-repeats count	Count of total subject repeats	Count of total put on subjects	Total accomplished credit	Calculated rate
1.	0	0	76	213	2,8026
2.	1	1	16	74	4,6250
3.	0	0	81	223	2,7531
4.	3	5	26	85	3,2692
5.	0	0	60	191	3,1833
6.	0	0	55	195	3,5455
7.	2	3	36	83	2,3056
8.	1	1	33	118	3,5758
9.	0	0	13	58	4,4615
10.	1	1	26	120	4,6154

Table 5: The subject-repeat statistic of the randomly chosen students

Number of the student	Different subject-repeats count	Count of total subject repeats	Count of total put on subjects	Total accomplished credit	Calculated rate
1.	2	2	41	140	3,4146
2.	6	7	76	183	2,4079
3.	7	7	64	157	2,4531
4.	13	17	80	178	2,2250
5.	18	30	98	176	1,7959
6.	4	4	40	130	3,2500
7.	3	3	65	168	2,5846
8.	0	0	13	58	4,4615
9.	1	1	32	116	3,6250
10.	0	0	11	53	4,8182

In the last column the calculated average came from the quotient of the total accomplished credit and the total put on subject count.

We made all calculation with the built in Analysis ToolPak of Microsoft Excel 2010 and 2013. You can install this addon in the File/Set up/Addons menu on the bottom of the opening window, clicking to the Jump button. After that you

have to choose on the Data tab Data analyzing menu item[4].

III. RESULTS AND DISCUSSION

At the beginning the averages were compared (figure 1). On the first chart we can see, that the competition participant students' entire three cases that means average grade, grade of Programming I. and Analysis I, are much better, than that of the other ten randomly chosen students.

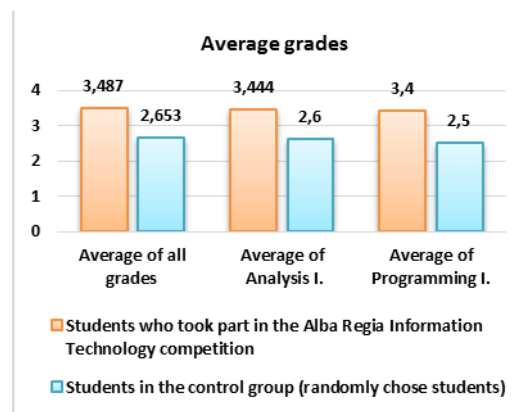


Figure 1: The compare of the averages

It seems to be nice, that the averages indicated differences in the results of the two student-groups, but we decided to make statistical test to support our hypothesis.

First we made an F-test to see, if the variances are equal or not in the two group. It is important, because according to the result of F-tests we had to choose the corresponding T-test. The result was at the total study average of the two groups: $F=2,129$ and $F_{critical}=3,178$. By the Analysis I. exam mark came $F=1,636$ and $F_{critical}=3,229$. And at last the Programming I. subject exam mark showed $F=1,866$ and $F_{critical}=3,178$. We can mark in all three cases, that F is smaller than $F_{critical}$. It signs that the variances are equal. We had to choose T-test for further calculations in the case if the variances are equal.

For null hypothesis we chose in all three cases, that there is no significant difference between the results of the two groups of students.

Let's see the results of the two sampled T-test. In the first case (study average) the result is: $t=2,450$ and $t_{two-edged critical}=2,100$, at Analysis I. exam mark T-test gave us $t=1,668$ and $t_{two-edged critical}=2,109$ and at last Programming I. exam marks T-test result: $t=2,377$ and $t_{two-edged critical}=2,100$.

Analyzing the numbers, we can realize, that only by Analysis I. the absolute value of t is smaller than $t_{two-edged critical}$. So in the subject Analysis I. there is no significant difference between the two groups at 95% significance level. We kept our null hypothesis in this case[3].

In the other two cases we had to drop our null hypothesis. We are concluding from these statistical results, that those students, who earlier took part in

our competition, they have significantly better performance in the global average and Programming I. subject at the Óbuda University in Székesfehérvár.

Now we are concentrating only for the earlier on our competition part taken people. We are interested in, if there is any connection between the result on the competition and the further study-outcome.

We made correlation analysis. The correlation coefficient value was by examining the global study averages $r = 0,215$. Count of the items was 10, so $r_{critical}$ was 0,632.

The absolute value of r is smaller than $r_{critical}$. That's why the hypothesized association rejected[3].

The same result was found by analyzing the Programming I. and Analysis I. exam results. The small difference is only, that by the Analysis I. marks we have a total item count nine, so $r_{critical}$ would be 0,666. By Analysis I. is the value of linear correlation coefficient $r = 0,323$, and by Programming I. r is 0,453.

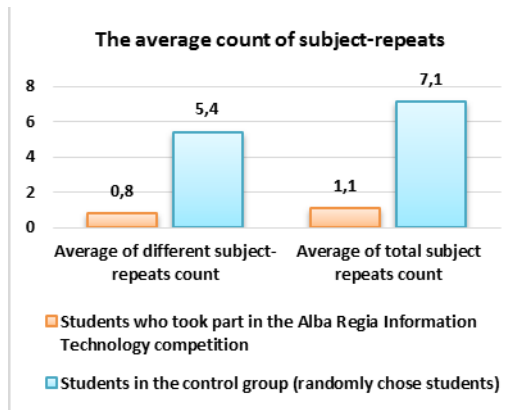


Figure 2: The comparison of count of the subject-repeats

After looking at figure 2 we can unambiguously say that our competing students are much better in proceeding of studies. They had to repeat a subject more rarely.

For the further evaluation of the data base we made a calculated rate of the division of the next two values: total accomplished credits and the count of total put on subjects. With this rate we would like to make more expressive how successful a student is. Here we can found only a little difference between the two averages. The average of the competing students for this calculated rate is 3,513, and the control group's average is 3,103.

Examined how many subjects need a student to repeat, we can realize the next. We made an F-test and T-test with the data of the different repeated subject counts. The variances were found equal, because $F = 0,030$ and $F_{critical} = 0,314$.

The T-test ended with the values $t = -2,416$ and $t_{two-edged critical} = 2,100$. The absolute value of t is not smaller than $t_{critical two-edged}$. That means that we have to drop our null hypothesis. We can say

that there is significant difference between the two groups of students by the count of repeated subjects.

Finally the calculated rates were investigated too. The result of the F-test displayed different variances. With the corresponding T-test we got the value for $t = 1,008$ and for $t_{two-edged critical} = 2,109$. The calculated rate is not significant different by the two groups.

IV. CONCLUSION

In this paper we compared two groups of students, each containing 10 persons. One group was the students, who took part earlier in the final of competition at our university, and the other group contained randomly picked out ten students from our study database. We searched the answer, if the teaching of applied information technology is able to develop general information technology capabilities, and the applied information skills are correlating with the mentality of this subject. Further we had a question if this competition is mending the enrollment of our faculty, and how. We displayed, that statistically the Programming I. exam results and the global study average was better by the competing students then the randomly chosen people. The applied information technology skill is thought can be capable for the selection. The information technology competition unambiguously mends the enrollment of our faculty. Further task can be, interpreting, if the competition is attractive for the students.

Interesting is, that by Analysis I. we didn't found significant difference by the two groups in contrast of our prejudice. It raises questions concerning the connection between analysis and programming teaching.

There is no correlation between the result of the competition and the further study outcome. Our opinion is that main factor is getting into the final competition. The study progress is significantly much better by the competing students. So we displayed, that the earlier competing students are going to learn better in the high education. But this is not valid for every subject.

What can be the reason? Our opinion is that the students, who were concerned with information technology in their free time, are more motivated.

The results of the statistical tests strengthened our hypothesis. But we have to handle them with criticism, because the small count of people in the groups.

This outcome of the examination is satisfying for us. It's a good idea and worth to organize this competition, because it gives a plus to the participant students. We can say decidedly that this year we'll organize this competition for secondary school students at Óbuda University Alba Regia Technical Faculty again, and we are waiting students too.

V. ACKNOWLEDGMENT

We acknowledge the financial support of this work by the Hungarian State and the European Community under the TÁMOP-4.2.2.B-15/1/KONV-2015-0010 project entitled "Development of Alba Regia Technical Faculty's Scientific Workshops".

VI. REFERENCES

- [1] "Alba Regia Information Technology Competition" [Online]. Available: <http://infoverseny.amk.uni-obuda.hu>.
- [2] A. Dávid, B. Maróti, Á. Burián, and É. Hajnal, "Alba Regia Competition of Applied Informatics as a Special E-Learning Field," in *7th International Symposium on Applied Informatics and Related Areas*, Székesfehérvár: Óbudai Egyetem, 2012, pp. 57–60.
- [3] P. G. Hoel and others, "Introduction to mathematical statistics.," *Introduction to mathematical statistics.*, no. 2nd Ed, 1954.
- [4] K. Berk and P. Carey, *Data Analysis with Microsoft Excel: Updated for Office 2007*. Cengage Learning, 2009.

Data collection and analysis of heart rate variability with the emWave2: a case study

Ildiko Nyireti¹, Melitta Szanto¹ and Jozsef Halasz^{1,2}

1. Obuda University, AMK, Szekesfehervar, Hungary

2. Vadaskert Child Psychiatry Hospital, Budapest, Hungary

nyireti.ildiko@gmail.com

szanto.melitta@gmail.com

halasz.jozsef@amk.uni-obuda.hu

Abstract— Heart rate variability is an important measure of basic cardiovascular functions, and extensively used in stress research. A portable system was introduced into the market for biofeedback applications, detecting interbeat intervals and coherence in various situations and training conditions. The aim of the present study was to describe first-order preliminary data in a case study with the help of the emWave2 system, and to test the possibility for applications in real-life situations with natural context. Additional to the first-order analysis, an advanced analysis with the help of Kubios software was also introduced, and future possible applications are also outlined.

Keywords: Cardiovascular, Heart Rate Variability, Natural Context, Stress.

increase in the number of commercially available systems [6] that can be used in everyday situations. E.g., some systems provide the heart rate information during jogging or different physical activities, some containing an additional activity tracker [also calculating calorimetric changes], that might be used for monitoring diurnal activity. Another direction is the application of different biofeedback systems, like HeartMath applications [7], where targeting stress and maintaining resilience is the main question [8]. The emWave2 system has been used in recent studies, improved coherence measures were associated with better physical and emotional parameters after regular emWave2 usage in populations with marked emotional burden, major depression and hypertension [9-14].

I. INTRODUCTION

A. Heart rate variability (HRV)

Heart rate is a nonstationary signal, dynamically changes during different physical or psychological stress situations. Original measures of basic cardiovascular parameters like heart rate and blood pressure have been supported with HRV analysis (variations in the interbeat intervals) [1-3]. HRV is a complex measure of cardiac functions and the state of the autonomic nervous system responsible for regulating cardiac activity. HRV has a major impact assessing sympathetic and parasympathetic elements, with a major effect of different emotional states, cognitive load or cardiovascular diseases [3]. E.g., based on the Framingham Study, repeated analyses of cardiovascular mortality were associated with decreased heart rate variability [4], and multiple associations were also described with other diseases as well [5].

Major neurophysiological laboratories use advanced systems to analyze heart rate variability during different laboratory conditions; but “more natural” portable systems (with less precision) have also major advantages in modeling real life situations. Technological development in the last 5 years created an unparalleled

B. Aims

The aims of the present experiments were twofold. *First*, the emWave2 system was used for detecting heart rate measures in a case study during different resting conditions (control condition, slow and fast music) and during moderate physical activity. Our hypotheses were the followings: compared to baseline resting conditions, slow and fast music might increase baseline heart rate in resting conditions, and physical activity would decrease interbeat intervals.

Second, our aim was to model first-level and advanced analysis of the heart rate on the present set of data.

These data served as preliminary data for a study in adolescents with externalization problems, the applicability of the system in different natural settings was tested. The detailed description of the above study was beyond the scope of the present paper [for background, see 15], the research entitled “Dimensional approach in externalization disorders” was approved by the Scientific and Research Ethical Committee of the Hungarian Medical Research Council (ETT-TUKÉB).

II. METHODS

The HeartMath emWave2 Portable Stress Relief system has a size of 85mm x 55mm x 14mm, handy and easy to

use. The setup has two types of sensors, both the ear-clip and the touch sensor (thumb or index) can be used in natural settings [7]. The system can be connected with a USB plug-in to a notebook, with a user friendly interface. Data collection can also be performed without the computer, and later data acquisition can be performed. The setup has two types of sensors. In the present set of experiments, a large number of situations were repeatedly tested, what can occur during regular daily activities (more than 20 conditions, e.g. breakfast, learning, reading, cooking etc.). Data collection lasted approximately for two days. Ear-clip and touch sensor data was also used, only the data for the touch sensor were shown. The test subject was a university student, female, 21 years old, one of the authors (I. Ny.). Her data are shown in the results section, in selected scenes.

The duration of the measurements was 180 seconds for different resting conditions. Slow (Titanic – My heart will go on, Celine Dion) and fast (Danza Kuduro, Don Omar) music condition was defined as resting condition and additional listening to the given music. In another setting, the effect of mild physical activity was tested (90 seconds). The system was able to provide sec by sec heart rate data and interbeat intervals, all the further steps were performed by the authors with the help of Statistica 7.0 and Kubios softwares.

Table 1. Heart rate data (per min) in different conditions in 30 sec packages

Heart rate		fast music	slow music	rest
first 30sec	average	69,7	75,9	70,5
	min	59,1	66,9	56,3
	max	82,2	93,6	86,7
second 30sec	average	75,0	81,7	69,1
	min	68,3	71,3	61,3
	max	92,7	97,7	86,8
third 30sec	average	74,4	72,8	68,8
	min	67,7	61,6	61,1
	max	83,2	79,6	79,8
fourth 30sec	average	72,5	73,5	70,8
	min	63,4	61,5	55,8
	max	84,7	80,8	85,2
fifth 30sec	average	75,7	75,0	68,7
	min	69,6	60,2	61,5
	max	82,4	90,6	78,9
sixth 30sec	average	74,2	66,6	68,1
	min	69,7	56,6	63,8
	max	80,3	71,7	73,3

Statistical analysis. As a case study was performed, only non-conventional measures could be used, with self-control. Statistica 7.0 was used for a GLM-based testing in the comparison of different resting states (resting vs. slow music vs. fast music), and the calculated (by emWave) 500 msec fluent heart rate data were compared in 30 sec packages. Neumann-Keuls post hoc comparisons were also run, $p < 0.05$ was considered as significant. In the comparison of interbeat intervals between resting state vs. moderate physical activity, the interbeat intervals were analyzed for 120 intervals in each

condition. In the latter context, frequency spectrum analysis and nonlinear analysis (Poincare plot) was applied with the help of the Kubios HRV 2.1 software (Department of Applied Physics, University of Eastern Finland).

III. RESULTS

In the first set of data, descriptive statistics of different resting conditions (rest [without music] vs. rest with slow music vs. rest with fast music) in 30 sec packages were presented (Table 1). The average, minimal and maximal heart rate data were outlined. In the first 30 sec, a significant difference occurred between groups ($F_{(2,180)}=23.45$, $p < 0.0001$), slow music was accompanied by significantly higher heart rate data compared to the resting condition (Fig. 1.). In the next four 30 sec packages, a significant increase was present in both slow and fast music conditions compared to heart rate values in resting condition (Fig. 1.; 2nd: $F_{(2,180)}=78.31$, $p < 0.0001$; 3rd: $F_{(2,180)}=26.75$, $p < 0.0001$; 4th: $F_{(2,180)}=5.94$, $p < 0.005$; 5th: $F_{(2,180)}=34.92$, $p < 0.0001$). In the last, 6th package, a significant ($F_{(2,180)}=135.40$, $p < 0.0001$) but differential effect was observed: while fast music was associated with a significant increase of the heart rate, slow music was associated with lower values than in the control resting condition (Fig. 1.).

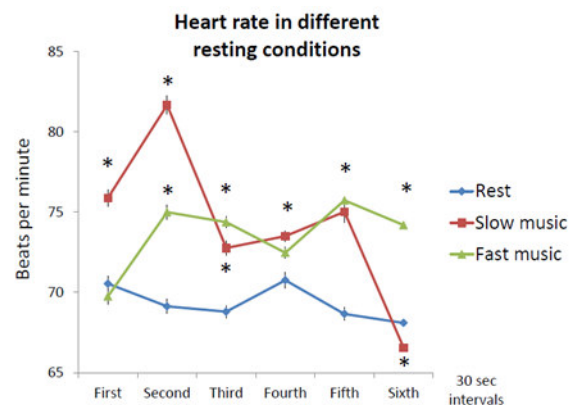


Figure 1. Heart rate data acquired with the emWave2 setup in three different conditions (averages for 30 sec, from 0.5 sec transformed data based on individual interbeat intervals). *, $p < 0.05$ from the corresponding resting condition.

In the second set of data, a different type of analysis was applied. A detailed histogram shows a marked difference between resting condition and moderate physical activity (Fig. 2.). The average heart rate data for resting condition was 69.1/min, while 89.4/min for moderate physical activity. The average interbeat interval was 876.9 ms and 688.0 ms for the corresponding groups. There was a significant decrease in the interbeat interval values ($F_{(1,238)}=210.05$, $p < 0.0001$) during physical activity compared to resting conditions.

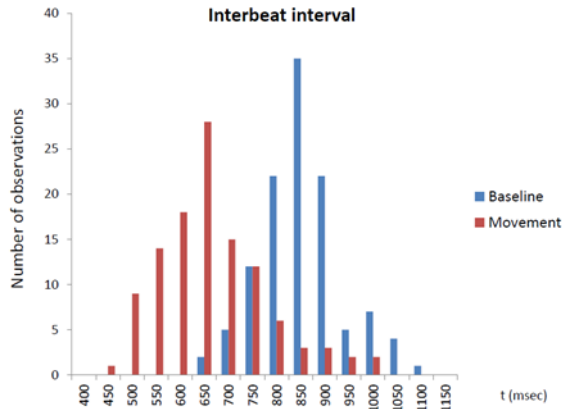


Figure 2. Interbeat interval data acquired with the emWave2 setup in two different conditions (Baseline – resting condition; Movement – mild physical activity). Interbeat interval categories were set in steps of 50 msec.

Additional to the first level analysis, an advanced

analysis was also performed with the second set of data with the help of the Kubios software. Sample report pages for resting baseline conditions (Fig. 3.) and moderate physical activity (Fig. 4.) were outlined. Major differences were described between the frequency-domain spectra and in the Poincare plot.

IV. DISCUSSION

The main results of the present case study were the followings. First, in our preliminary study, with the help of the emWave2 setup, heart rate variability measures could be reliably assessed. Second, subtle differences could be outlined in the above conditions, and these data could be used for orientation in our planned experiments in adolescents with externalization problems.

As it was mentioned in the Introduction, emWave2 is a user-friendly biofeedback instrument, what can be used in various natural conditions [7,8,12-14]. Recently, the measured improvement in coherence was not only associated with better cardiovascular parameters but also with significant increase in self-control, lower anxiety, and improved quality of life [9-14]. So far, we do not have information about the usage of the setup in

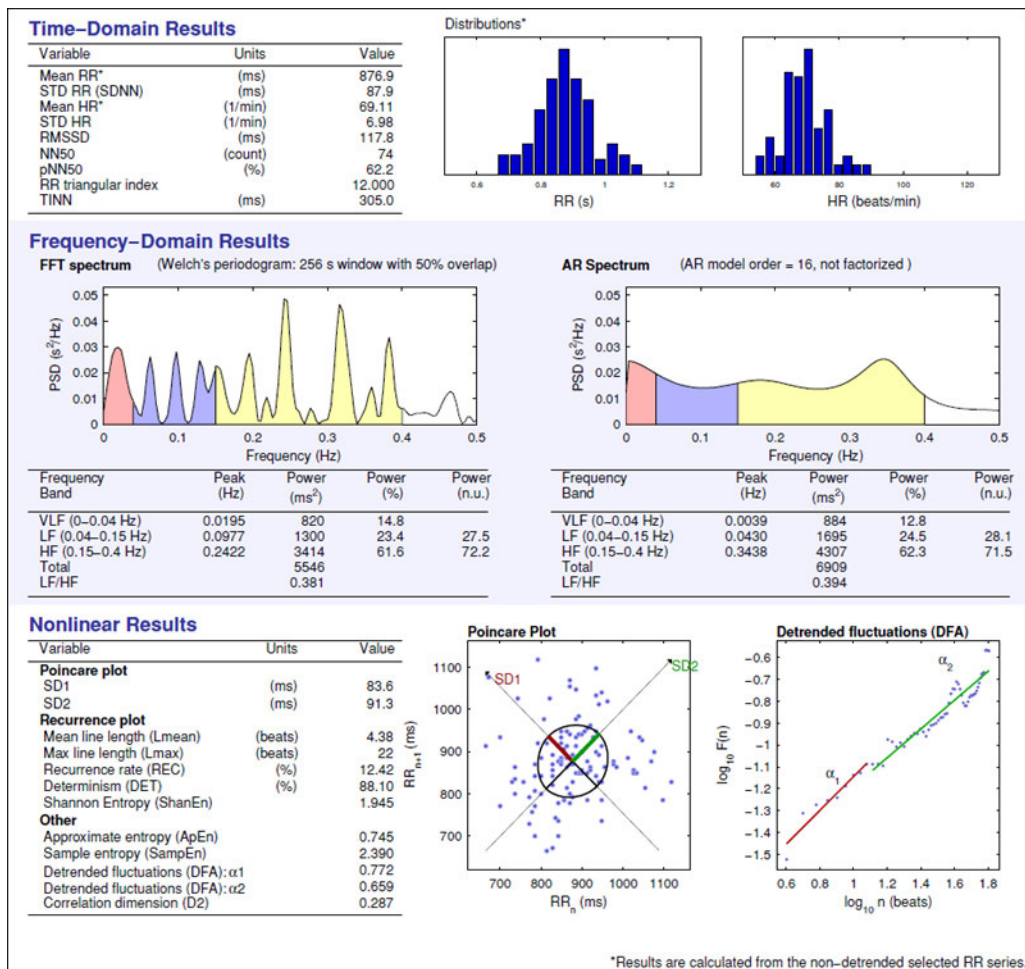


Figure 3. Time domain, frequency domain and nonlinear advanced analysis data from the Kubios software of 120 interbeat intervals in baseline resting condition (emWave2).

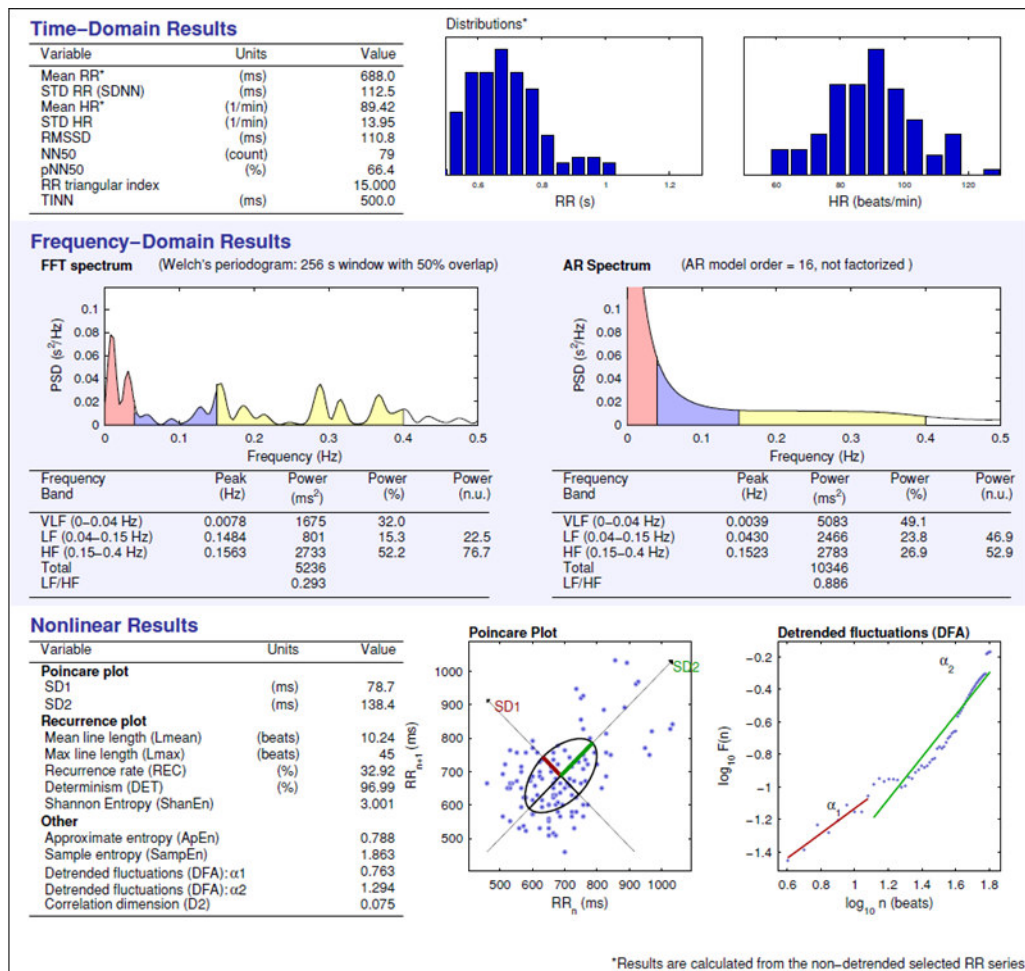


Figure 4. Time domain, frequency domain and nonlinear advanced analysis data from the Kubios software of 120 interbeat intervals during mild physical activity (emWave2).

adolescents in general or adolescent with externalization problems in particular. Nevertheless, in our preliminary case study the biofeedback effect was not used, the setup was only used for advanced data collection. As no “A” evidence exist for the successful treatment of adolescents with externalization problems [16], the future therapeutic application of a biofeedback system in anger control and improved self-management might be addressed in affected adolescents.

Heart rate related data in adolescents deliver way more message than vegetative control or stress related parameters. E.g., externalization problems might be characterized by alterations in the working of limbic brain areas, particularly the prefrontal cortex and the amygdala [17–20]. Damasio’s famous intention on the relationship between the somato-vegetative and cognitive functioning based on the working of the prefrontal cortex [21] might also be taken into account. According to the somatic marker hypothesis, emotion regulation and control is highly associated with monitoring of somato-vegetative signals, thus biofeedback training in developing individuals might be especially important. As the development of the prefrontal cortex is not finished during adolescence, biofeedback setups might also be important in the “somato-vegetative training” of the

prefrontal cortex in conditions where altered prefrontal functioning is associated with behavioral vulnerability [17].

In the present setting, only different levels in the analysis of heart rate and interbeat interval data were addressed. With the user-friendly emWave2 setup dynamic changes of the heart rate and oscillating spectral features easily could be assessed. Unfortunately, so far we do not have any information of child psychiatric studies targeting externalization problems in adolescents in a frame where advanced spectral and nonlinear analysis was performed in heart rate variability. Our preliminary data raise the possibility for analyzing emWave2 data with the advanced Kubios software. Detailed description of the parameters calculated within the software can be found e.g. in [5], and hopefully, behavior oriented procedures would also apply the methodology (and not just cardiology oriented research). Nevertheless, our hypotheses were fulfilled, and statistical differences could be obtained. In the first set of data, within the resting condition, rather low fluctuation in 30sec heart rate averages could be observed, while listening to slow or fast music qualitatively changed the pattern. Also, marked difference between resting and physical activity was prominent, and advanced data

analysis revealed significant differences, what might be associated with cardiovascular adaptive functions.

The limitations of the study were the followings. This was a case study for preliminary data collection, thus the application of the methodology was the main question, and the design could not be compared to the size of a controlled study. Another weakness of the study that additional to the internal control of the system, external control (e.g. application of a parallel system with strict ECG record) was not applied. In future studies, these issues also have to be addressed.

V. SUMMARY

In the present paper, a case study with preliminary data collection of heart rate variability with emWave2 portable setup was described. Numerous settings were tested, and among them, two were described in details, and corresponding data analysis with the advanced Kubios software was outlined. The present data were necessary to arrange further steps in our main research in adolescents with externalization problems, where the usage of the setup during different neuropsychological tests would create a so far unique combination in our specific area.

VI. ACKNOWLEDGMENT

The authors declare no conflict of interest. The work was supported by OTKA Grant No. 108336.

REFERENCES

- [1] H.M. Stauss, "Heart rate variability", *Am. J. Physiol., Reg., Integr. Comp. Physiol.*, vol. 285, pp. R927–R931, 2003.
- [2] G.G. Bernston, J.T. Bigger, D.L. Eckberg, P. Grossman, P.G. Kaufmann, M. Malik, H.N. Nagaraya, S.W. Porges, J.P. Saul, P.H. Stone and M.W. van der Molen, "Heart rate variability: Origins, methods and interpretive caveats", *Psychophysiology*, vol. 34, pp. 623–648, 1997.
- [3] R. McCraty, M. Atkinson, W.A. Tiller, G. Rein and A.D. Watkins, "The effects of emotions on short-term power spectrum analysis of heart rate variability", *Am. J. Cardiology*, vol. 76, pp. 1089–1093, 1995.
- [4] H. Tsuji, M.G. Larson, F.J. Venditti, E.S. Manders, J.C. Evans, C.L. Feldman and D. Levy, "Impact of reduced heart rate variability on risk for cardiac events", *Circulation*, vol. 94, pp. 2850–2855, 1996.
- [5] U.R. Acharya, K.P. Joseph, N. Kannathal, C.M. Lim and J.S. Suri, "Heart rate variability: a review", *Med. Bio. Eng. Comput.*, vol. 44, pp. 1031–1051, 2006.
- [6] J. Zheng, Y. Shen, Z. Zhang, T. Wu, G. Zhang and H. Lu, "Emerging wearable medical devices towards personalized healthcare", *BODYNETS*, pp. 427–431, 2013.
- [7] HeartMath, "emWave2 quick start guide for PC and Mac", *Quantum Intech*, pp. 1–20, 2014.
- [8] E.D. Lackey, "Self-regulation and heart rate variability coherence: promoting psychological resilience in healthcare leaders", Dissertation, *Benedictine University*, pp. 1–189, 2014.
- [9] J. Beckham, T.B. Greene and S. Meltzer-Brody, "A pilot study of heart rate variability biofeedback therapy in the treatment of perinatal depression on a specialized perinatal psychiatry inpatient unit", *Arch. Womens Ment. Health*, vol. 16, pp. 59–65, 2013.
- [10] C.M. Sarabia-Cobo, "Heart coherence: a new tool in the management of stress on professionals and family caregivers of patients with dementia", *Appl. Psychophysiol. Biofeedback*, DOI 10.1007/s10484. pp. 1–9, 2015.
- [11] F.J. Reyes, "Implementing heart rate variability biofeedback groups for veterans with posttraumatic stress disorder", *Biofeedback*, vol. 42, pp. 137–142, 2014.
- [12] S.D. Edwards, "Effects of biofeedback training on physiological coherence, health and spirituality perceptions", *AJPHRD*, vol. 20, pp. 500–510, 2014.
- [13] S.D. Edwards, "Evaluation of a heartmath workshop for physiological and psychological variables", *AJPHRD*, vol. 20, pp. 236–245, 2014.
- [14] S.D. Edwards, "Evaluation of heart rhythm coherence feedback training on physiological and psychological variables", *South African J. Psychol.*, vol. 44, pp. 73–82, 2014.
- [15] J. Halasz, N. Aspan, P. Vida and J. Gadoros, "Recognition of emotions in facial expressions- implications and limitations during antisocial development", 7th International Symposium on Applied Informatics and Related Areas, pp. 44–49, 2012.
- [16] National Institute for Health and Care Excellence, "Antisocial behaviour and conduct disorder in children and young people. Recognition, intervention and management", British Psychological Society and and The Royal College of Psychiatrists, *National Clinical Guideline Number 158*, pp. 1–468, 2013.
- [17] S.H. van Goozen, G. Fairchild, H. Snoek and G.T. Harold, "The evidence for a neurobiological model of childhood antisocial behavior", *Psychol. Bull.*, vol. 133, pp. 149–182, 2007.
- [18] A.A. Marsh, E.C. Finger, D.G. Mitchell, M.E. Reid, C. Sims, D.S. Cosson, K.E. Towbin, E. Leibenluft, D.S. Pine, and R.J. Blair, "Reduced amygdala response to fearful expressions in children and adolescents with callous-unemotional traits and disruptive behavior disorders", *Am. J. Psychiatry*, vol. 165, pp. 712–720, 2008.
- [19] A.P. Jones, K.R. Laurens, C.M. Herba and G.J. Barker, "Amygdala hypoactivity to fearful faces in boys with conduct problems and callous-unemotional traits", *Am. J. Psychiatry*, vol. 166, pp. 95–102, 2009.
- [20] R.J. Davidson, K.M. Putnam and C.L. Larson, "Dysfunction in the neural circuitry of emotion regulation--a possible prelude to violence", *Science*, vol. 289, pp. 591–594, 2000.
- [21] A. Bechara, H. Damasio, D. Tranel and A.R. Damasio, "Deciding advantageously before knowing the advantageous strategy", *Science*, vol. 275, pp. 1293–1295, 1997.

The Use of Self-organizing Maps (SOM) in Demographic Analysis

Nagy Balázs*, Majdán Márk*, Dr. Varga Valéria*, Dr. Nagyné Dr. Hajnal Éva*

* Óbuda University, Alba Regia Technical Faculty

Abstract—The present study investigates Hungary's higher education student numbers based on demographic as well as education system statistics. The following data sources were used: Central Statistical Office (Hungary), education information portal felvi.hu and Educational Authority (Hungary). The applied method, Kohonen's self-organizing map (SOM). In the interval between 2000 and 2012, the following statistics were used: total population, students in secondary education, students in higher education, those completing secondary education, fresh graduates and the number of students admitted to higher education; in the interval between 1982 and 2012 statistics of infant mortality and live births were used. We created two statistical scenarios to predict the number of years spent in higher education. After importing data into a SOM data structure and training the map, two clusters were recognizable, the separating line between the two being the year 2008. Low quantization and topographic error values indicate the high quality of data processing. After training the map, we made a prediction for the interval 2001-2030. Predicted values were retrieved after calculating the mean value of neurons linked to values in the weight vector matrix. The results suggest that by 2030, higher education student numbers will decrease by a percentage 30 to 35, as compared to the numbers in 2011.

I. INTRODUCTION

The research and development (commonly abbreviated as R&D) sector is a noteworthy factor in an economy due to its role in a country's competitiveness. The sector can be described with many indicators, however, higher education numbers are of special importance, as universities provide a new generation of research staff. The further advantage of this approach is that higher education statistics are widely available.

Higher education student numbers are determined by demographics and legal changes in education policy. In Hungary, decreasing demographic tendencies result in decreasing student numbers. A suitable method is necessary in order to be able to estimate the exact tendency of this change.

Although not primarily used for prediction and estimation, the method we chose is widely used for data analysis. In the present study we give an overview of how SOMs may be used for prediction and estimation, through establishing and evaluating a statistical model capable of estimating higher education student numbers from demographic data with good approximation.

Higher education institutions, apart from providing academic education, participate in various R&D activities. The border between universities as degree-giving

institutions and research institutes is hard to tell. Students wishing to obtain a PhD degree are required to have done a certain amount of research work, from which further publications and experimental results originate, even if students themselves do not stay on a researcher's career path.

In the past twenty years it has become common to build so-called knowledge centres, where representative offices of companies, state-funded or privately owned research institutes work on joint projects at the same place. This type of research cooperation is not fully developed in Hungary.

The structure and quality of Hungary's higher education is essentially determined by education policy reforms, such as the introduction of the so-called two-level exam system, replacing university entrance exams.

The Bologna Process, introduced in Hungary in 2006, has a significant effect on the structure and quality of academic education. The formerly five-year-long university courses (apart from a few) were divided into a three-year-long bachelor (BSc/BA) and a two-year-old master's (MSc/MA) course. Similarly the variety of university majors were also reduced: compared to 400 majors in 2005, applicants could only choose from 108 majors in 2006.

The decade between 2000 and 2010 was characterized by an increase in higher education student numbers. Based on higher education students per total population ratio, more people enrolled in university courses than ever before. Some distorting factors, however, should be taken into account. Applicants for master's courses after 2006 should also be counted in statistics of university applicant numbers.

Despite the above mentioned reforms, the most significant factor that affected higher education student numbers was demographics. In an analysis of long time series statistics, decreasing population and decreasing student numbers are clearly to be seen. These tendencies are only slightly changed by fluctuations caused by education reforms. Similarly to how age groups (demographic cohorts) can be tracked through the statistics of a population, a relatively good estimation can also be given for the ratio of students being admitted to higher education.

II. DATA AND METHODS

A. Data selection

Data from various sources were used in the present study. The main data sources were Hungary's Central Statistical Office (KSH) [1][2], education portals felvi.hu [3] and Education Authority [4]. In total, 152 data values

were used as well as data derived from the original data set.

The following from KSH were used: total population, mortality and birth rates, education statistics. Demographic data included birth rate (1982-2012; a total of 31 data values), total population (2000-2012; 13 data values) and infant mortality (1982-2012; 31 data values).

Education statistics included the following time series data: students in secondary education (2000-2012; 13 data values), students in higher education (2000-2012; 13 data values), students completing secondary education (2000-2012; 13 data values), higher education students receiving a degree (2000-2012; 13 data values).

Higher education data were retrieved from the website felvi.hu and the Education Authority. From felvi.hu's statistics, the following data sets were used: statistics of recent years, number of students applying for a university and number of those admitted to a university, from which we retrieved the number of students admitted to a university (2001-2012, 12 data values).

We obtained the number of first-year university students using Education Authority's higher education data (2000-2012, 13 data values).

B. An overview of the model

We created the prediction using artificial neural networks. Artificial neural networks are widely used to solve mathematical problems, also in statistical data processing, especially in the case of large and complex data sets. Pattern recognition neural networks are also used in face recognition [5][6][7] and optical character recognition, but also in robot navigation [8] (to avoid obstacles in a robot's path).

The method of choice is Kohonen's self-organizing map, abbreviated as SOM, created by Teuvo Kohonen [9].

This type of artificial neural network belongs to neural networks with unsupervised learning, this means that only an input is given without any reward signal. In a neural network, neuron connections are structured in a two-layer topology that consists of an input and an output layer. Some neural networks have further layers, also known as hidden layers, but Kohonen SOMs only have the two basic neuron layers. The output layer is two-dimensional and map-like, hence its name. Neurons in the output layer are interconnected and are located in a hexagonal or rectangular lattice.

Input signals activate neurons in certain areas of the output layer in case of similar activation patterns. It is important to know how neurons are interconnected in order to be able to optimise the activation patterns. The optimisation process creates a topological map of the input signals and transforms similarities between input signals into a neighbourhood indicator between neurons. This results in the map being able to ignore disturbing factors to take into account only the most predominant similarities (connections). In short, Kohonen's SOM compares input signals (data values), then uses these numeric values to calculate the neighbourhood of output neurons.

Each neuron's position in the lattice can be described with a weight vector assigned to it. Similarly to biological neurons, artificial neurons are also in a network of synaptic connections. This similarity can even be extended to the two basic types of synapses, inhibitory

and excitatory, the principle of which also exists in artificial neural networks.

Kohonen's SOM algorithm can be broken down into the following main steps. Let x_k denote a training sample of n dimensions and w_{ij} the neuron to be selected. The distance between neurons is defined by the neighbourhood function $h(w_{ij}, w_{mn})$. The neighbourhood function returns high values if neurons are close to each other in the output space. The selected winner neuron from the input data is w_{winner} (best-matching unit, abbreviated as BMU).

The following steps are completed for each input data value:

1. The distance between the input data value and all the other neurons is calculated one by one.

$$d_{ij} = \|x_k - w_{ij}\|$$

2. The closest neuron (the one with the smallest distance) is the winner neuron.

$$w_{nyertes}(w_{ij} : d_{ij} = \min(d_{mn}))$$

3. After selecting the winner neuron, the values of all the other neurons are updated.

$$w_{ij} = w_{ij} + \alpha \cdot h(w_{nyertes}, w_{ij}) \cdot \|x_k - w_{ij}\|$$

4. The process restarts after the third step and the distance between each input value and the winner neuron is calculated. The process is repeated until it reaches a previously set condition such as the number of iterations.

The distance between weight vectors can be calculated with different mathematical methods, but Euclidean metric is the most often used and we also used this type of metric.

The SOM updates weight vector values one by one until they get close to the BMU. This type of algorithm is called an adaptive algorithm. Learning rate (α) is a numerical value that determines how close neurons can get to the winner neuron in each learning epoch. Training can be iterative, when the entire data set is processed through iterations (steps) in the above mentioned way.

Another approach is batch training when the entire data set is partitioned into Voronoi cells based on similarities between data values. Voronoi cells are grouped around a neuron. By default, most self-organizing maps use iterative training.

After training the SOM, patterns close to each other in the input space are close to each other in the output space and vice versa. Thus, the SOM training algorithm preserves topology (neighbourhood values). Regardless of the number of training epochs, there is always a difference between the input value and the value assigned to a neuron. This is known as quantization error. It is worth to note that SOM training algorithm is similar to vector quantization (VQ), a type of unsupervised learning used in digital signal processing and data processing.

The output space can be described with topographic error. This number is calculated by measuring the distance of an input vector's BMU and its second BMU.

Aside from visualizing the output layer itself, a SOM also visualizes the U-matrix, which denotes with colour tones the Euclidean distance of weight vectors of neighbouring neurons. It provides useful information related to the entire data set. Clusters are denoted by light colours, whereas borders between clusters are denoted by dark colours.

SOMs, as clustering methods, are widely used for processing economic data. Pablo-Martí [10] demonstrated the regional distribution of economic activities in the Spanish economy; Collan, Eklund and Back [11] used a SOM to create a map that groups the countries of the world by their economic development. The above examples mostly take advantage of the two-dimensional, map-like visualization, as this type of visualization can be projected onto a geographical map.

Several examples can be found in bibliography for SOMs used as statistical estimation methods [12], often in various areas of sciences. Al-Shayea et al. [13] created a data prediction model from the financial data of sixty-six Spanish banks that is capable of predicting bank insolvency. Ultsch and Röske [14] created a sea level prediction model from historic meteorological data using a SOM.

C. Creating the model

We used Microsoft Office 2010 Excel to make estimations from the data available. Thus, we obtained the number of the 18-year-old age group by shifting the number of live births (2000-2030, 31 data values), the number of higher education students minus the number of those applying for a university (2000-2012, 13 data values) and the number of years spent in higher education (2001-2012, 12 data values). We obtained this by taking the number of students in higher education minus the number of those applying for a university and dividing it by the number of freshly graduated students.

The majority of estimations were conducted with Microsoft Excel. As for the final prediction, however, we used MATLAB and the MATLAB function library SOM Toolbox. We used an Excel estimation (2013-2030) to obtain the number of those receiving a university degree (18 data values), the number of years spent in higher education in two scenarios. One scenario is that the number of years spent in higher education follows an exponential trendline. The second scenario is that the number of years is at a constant average. Thus, we obtained the number of students minus those applying for a university (18 data values), the number of students * years and the number of students.

Sixty data values were used to conduct the MATLAB prediction. These were the following: the number of students * years spent in higher education, the number of the 18-year-old age group, the number of years spent in higher education, the number of those receiving a university degree and the number of students completing secondary education. In order to reveal connections between data values we used SOM Toolbox, which visualizes clusters and gradients to analyse data connections.

We used various settings to train the SOM. All the 60 available data values were imported into a sample matrix of the dimensions 12×5. A data struct 'D' was created from the sample matrix, then the data struct was

normalized, using linear 'range' type normalization. 'Range' type means that data values fall into an interval between 0 and 1, proportional to the biggest value of the data series. After this, we trained the SOM from the normalized data using Euclidean metric to measure distances. Euclidean metric is based on Pythagorean theorem and bibliography suggests it is a robust metric to be used in this case.

To create the prediction model, we used a SOM similar to the previous one, but in this case with data in the interval between 2001 and 2030. The model was created in two scenarios. In the first one we supposed that the number of years spent in higher education follows an exponential trendline, whereas in the second one the number of years was a constant mean value. After training the map, the four neurons (BMUs) closest to each weight vector were determined. Then BMUs of each row in the weight vector matrix were placed in the variable 'index'. This variable helped us to find values in the weight vector matrix associated to the index. The next step was creating a matrix in order to store the mean value of weight vectors. The final step was data denormalization, after which we obtained the results.

III. RESULTS AND DISCUSSION

In the first step, we provided a common number format for all the data series available.

We investigated the number of participants in secondary education, those enrolled in higher education, those completing secondary education and those receiving a degree.

The number of those completing secondary education, those enrolled in secondary education and those receiving a degree began stagnating in 2006 and after 2008 a drop in their numbers can be observed. However, in the years after 2008, the numbers reflect the values before the rapid decrease and are constant apart from minor fluctuations.

In the higher education, however, different tendencies can be observed. One of the reasons inducing an increase in student numbers was an 1985 law on higher education and its modification imposed in 1990, as these caused large numbers of students to be admitted to higher education. The subsequent laws on higher education also increased student numbers.

The increase reached its peak point in 2006, when the new law on higher education was imposed, thus introducing the Bologna Process into Hungary's education system. This caused a drop in student numbers in Hungary. One reason for this was that by the new law, students had to pay for subjects not included in the university curriculum. Another reason is that stricter regulation was introduced concerning language proficiency as prerequisite of graduation and also for university admission.

The most significant decrease starts from 2008. This is partly due to the effects of the economic crisis and also due to growing tuition fees.

Our aim was to give an estimation for higher education student numbers twenty years ahead. As the first step, we determined the number of the 18-year-old age group by shifting birth numbers 18 years ahead in time. As infant mortality is low, it is not a significant factor that would distort our estimation.

Knowing the number of the 18-year-old age group, the second step was determining what percentage of this age group completed secondary education and was enrolled in higher education. In the last decade, the number of students completing secondary education stagnated at approx. 71.13%. We then applied this ratio to the number of the 18-year-old age group.

After this we determined the number of students in higher education by subtracting the number of recently admitted students from the total higher education student number. We took into account not just bachelor students, but also those enrolled in master's and PhD courses. We then divided the resulting value by the number of those receiving a degree, thereby giving an estimated mean value of 5.5 years for the number of years spent in higher education.

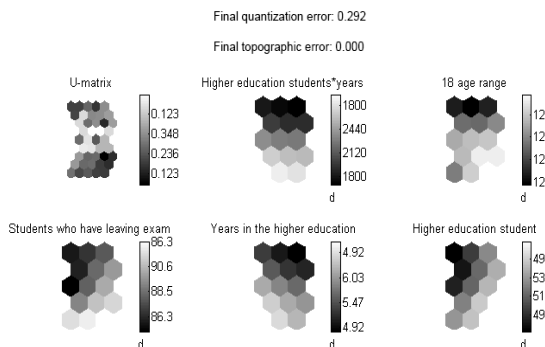


Fig. 1: A SOM after training with the higher education sample data. Five different variables were used to train the SOM. The U-matrix shows two, clearly distinguishable clusters, the separating value between the two clusters is the year 2008.

In order to estimate the number of higher education students we had to determine the number of those receiving a degree and those having spent at least one semester in higher education. We created a diagram from these data and drew a linear trendline to determine the number of those receiving a degree twenty years ahead. In the period between 2005 and 2008, the number of those receiving a degree showed a decreasing tendency, followed by a temporary increase between 2008 and 2011, when it started to decrease again. If this decreasing tendency continues, then student numbers are likely to drop to 40.000 in 2030. However, this is still a rough estimation, which needs to be improved via the SOM method.

Now all the necessary data were available to us to train a SOM and use it for prediction. In order to create the neural network we used the number of students*years spent in higher education, the number of the 18-year-old age group, the number of years, the number of those receiving a degree and those completing secondary education, all in the period between 2001-2011.

The map was created using 55 data values that were imported into a MATLAB variable. After creating the map, we visualized the results. At this stage there was no need to make any distinction between the first and the second scenarios, as in the period between 2001-2011 all the data values were available. Two clearly distinguishable clusters were present in the visualization, with the year 2008 separating them.

Quantization error shows the mean distance between input vectors and their BMUs. Topographic error shows

the percentage of data values where related BMUs are not neighbours to each other. In this case both error values were considerably low. Quantization error equals to 29%, whereas topographic error has a value of 0 %, which means related BMUs are in each other's neighbourhood.

We then estimated student numbers with a linear trendline from 2011 to 2030.

To estimate the student numbers we used two different methods. The first one was drawing a linear trendline with Excel.

Student numbers in 2030 were given using the equation of the Excel trendline. According to the results, student numbers are likely to decrease by 50% from present day to 2030. However, a linear trendline does not fit perfectly to the data series, as it would drop into the negative domain.

To give a more exact estimation for the one conducted with a linear trendline, we also made a prediction using SOM. As a first step, we multiplied the results from linear trendline estimation by the mean value of years spent in higher education (first scenario), then by the number of years from exponential trendline (second scenario). The two variants are largely similar, thus we expected similar results.

Now we had all the necessary data to conduct the SOM prediction. As with the data in the period between 2001 and 2011, we imported data in the period between 2001 and 2030 into MATLAB. After data normalization, the map was trained.

In this case, there are several data values where related BMUs are not in each other's neighbourhood. The bigger the topographic error value is, the more complex the output space is. Cluster patterns and topology itself in these cases are more complex. A high number suggests errors in the training process. Quantization error had a value of 9 and 10 %, respectively, topographic error had a value of 0 and 3 %, respectively, thus we can say the model proved to be reliable.

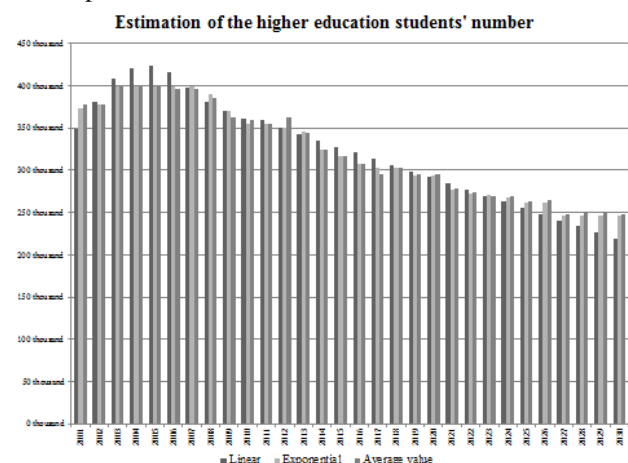


Fig. 2: Estimation of student numbers with linear trendline and with SOM (two different scenarios).

After creating the map we obtained the four neurons closest to each weight vector and looked up the values associated to them in the weight vector matrix. We obtained prediction results by calculating the mean of

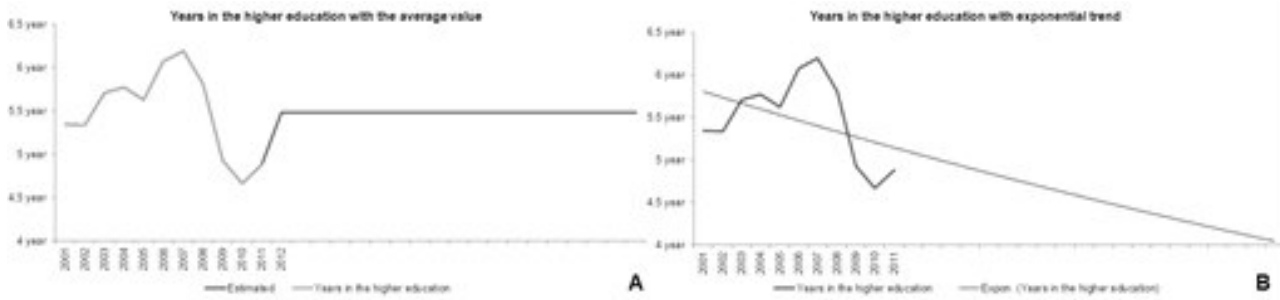


Fig. 3: The estimation of the years spent in the higher education. For the estimation data were used between 2001 and 2011. As seen on the graph the values changed between 4.5 and 6.5 years. A: On the graph the data was estimated with the average values. B: The exponential trend estimation shows that the number in 2030 might be between 4 and 4.5.

these values. Some recurrent values were also obtained because the method is extrema sensitive. This means that in some cases the same neurons are found to be the closest neurons, thus returns its values.

According to the SOM estimation, a decrease of 30% in student numbers is likely to occur in the first scenarios, whereas the second scenario predicts a decrease of 35% as compared to the situation in 2011. The linear trendline estimation thus gives a rough estimation as compared to the SOM prediction. Results only suggest one possible outcome, but many other should also be taken into account, such as students working and studying abroad. A larger amount of data is also necessary to validate the method.

IV. CONCLUSION

In the present study we tried to demonstrate the use of Kohonen's self-organizing map in the analysis of economic data through a case study. We estimated the number of higher education students twenty years ahead using SOM. According to data processing quality indicators, such as quantization error and topographic error, the model gives a good approximation of the observed tendencies. The present data quantity is not enough to test the validity of the method, but the method should be applied in case of larger quantities of historic time series data.

The decrease in student numbers are primarily caused by demographic tendencies, however, these were temporarily compensated by legal changes in the higher education. The tendency has been becoming more prevalent since 2008. Higher education produces researchers; significant R&D and innovation takes place in the sector itself, thus changes to higher education strongly determine the future situation of R&D as a sector of economy. Our study provides a basis for the research of these economic impacts.

In the future we would like to improve the prediction model based on several different scenarios. It should also be examined how the present rate of decrease of student numbers will affect the numbers of researchers and also the number of scientific publications and patents. In another scenario, an estimation should be given of the amount of financial assets necessary in order to maintain the present-day level of R&D activity while the number of research staff is decreasing.

ACKNOWLEDGMENT

We acknowledge the financial support of this work by the Hungarian State and the European Community under the TÁMOP-4.2.2.B-15/1/KONV-2015-0010 project entitled "Development of Alba Regia Technical Faculty's Scientific Workshops".

REFERENCES

- [1] K. S. Hivatal, "Népesség, népmozgalom 1949-" 2013. [Online]. https://www.ksh.hu/docs/hun/xstadat/xstadat_hosszu/h_wdsd001a.html.
- [2] K. S. Hivatal, "Oktatás (1960-)," 2013. [Online]. https://www.ksh.hu/docs/hun/xstadat/xstadat_hosszu/h_wdsi001b.html.
- [3] Felvi.hu, "Elmúlt évek statisztikái." 2013. [Online]. http://www.felvi.hu/felveteli/ponthatarok_rangsorok/elmult_evek.
- [4] Oktatási Hivatal, "Kezdő évfolyamos hallgatók száma képzési szintenként," 2013. [Online]. http://www.oktatas.hu/felsooktatas/felsooktatasi_statistikak/!DA RI_FelsooktStat/fir/fir_stat2010/stat2010_113.xls.
- [5] J. L. Alba, A. Pujol, and J. J. Villanueva, "Separating geometry from texture to improve face analysis," in IEEE International Conference on Image Processing, 2001, vol. 2, pp. 673–676.
- [6] C. H. Lee, D. S. Seong, and K. H. Park, "Face recognition using self-organizing map," J. Korea Inf. Sci. Soc., vol. 20, no. 11, pp. 1730–1738, Nov. 1993.
- [7] Z. Q. Liu, "Adaptive subspace self-organizing map and its applications in face recognition," International Journal of Image and Graphics Vol. 2, P. A.; Garrido Frenich, A.; Torres, J. A.; Pulido Bosch, A., no. 4, pp. 519–540, 2002.
- [8] I. K., N. S., and U. T., "A self-organizing map based navigation system for an underwater robot," in 2004 IEEE International Conference on Robotics and Automation IEEE Vol. 5, 2004, p. 5726.
- [9] T. Kohonen, "Self-organizing formation of topologically correct feature maps," Biol. Cyb., vol. 43, no. 1, pp. 59–69, 1982.
- [10] J. M. Pablo-Martí, F.; Arauzo-Carod, "Spatial distribution of economic activities: an empirical approach using self-organizing maps," in Rethinking the economic region. New possibilities of regional analysis from data at small scale, E. F.-V. F. Rubiera-Morollón, Ed. Springer, pp. 1–41.
- [11] M. Collan, T. Eklund, and B. Back, "Using the Self-Organizing Map to Visualize and Explore Socio-Economic Development," EBS Rev. 22(1), 6-15, vol. 22, no. 22, 2007.
- [12] C. Science, M. Po, T. Honkela, and T. Kohonen, Bibliography of self-organizing map (som) papers: 2002–2005 addendum. 2009, pp. 2002–2005.
- [13] Q. K. Al-shayea, G. A. El-refae, and S. F. El-itter, "Neural Networks in Bank Insolvency Prediction," IJCSNS International Journal of Computer Sciences and Network Security VOL.10 No.5, vol. 10, no. 5, pp. 240–245, 2010.
- [14] A. Ultsch and F. Röske, "Self-organizing feature maps predicting sea levels," Self-Organizing Feature Maps Predicting Sea Levels, Inf. Sci. 144/Elsevier, pp 91 - 125, vol. 144, pp. 91–125, 2002.

Quality Assessment System in Obuda University Alba Regia Technical Faculty

Margit Horosz-Gulyás, PhD

Alba Regia Technical Faculty of Óbuda University, Székesfehérvár, Hungary
horoszne.margit@amk.uni-obuda.hu

Abstract—What is important for the employee in its workplace (especially in higher education)? Why should the management know what motivates the staff? What is the role of the students? These are only a few questions, I looking for the answers. Due to a project a new assessment system was made based on the motivation of the employees.

Keywords: Quality, Employee, Motivation, Students

I. INTRODUCTION

We work in surroundings that our colleagues of thirty years ago would not recognize. Higher education has become part of a global shift to a new way of creating and using knowledge. The new way is focused on solving problems and be sensitive to customer needs. It strives for quantity as well as quality. It cuts across disciplinary boundaries.

In knowledge-based economies, governments see universities as engines for social change and the expansion of prosperity. University teachers have accordingly found themselves working harder and at the same time being required to be more businesslike and more accountable. The pleasures of the academic life have dwindled for many university teachers. They are unimpressed especially by the administrative effort associated with quality assurance and accountability. It uses up time and energy that could be focused on the core business of research and teaching [1].

Practically everybody in the academic community gets assessed these days, and practically everybody assesses somebody else. Students, of course, come in for a heavy dose of assessment, first from admissions offices, later from the professors who teach their classes, and increasingly from administrators complying with state accountability requirements. Students are also active participants in the assessment business, with end-of-course evaluations that are widely used by colleges and universities and various forms of web-based assessments of professors. The whole institution is regularly assessed in a highly detailed fashion by external accrediting teams made up of faculty and administrators from other institutions.

As commonly used today, the term assessment can refer to two different activities: the gathering of information (measurement) and the use of that information for institutional and individual improvement (evaluation) [2].

To manage change, universities and also other organizations must have employees committed to the demand of rapid change and as such committed employees are the source of competitive advantage. The concept of human capital and knowledge management is

that people possess skills, experience and knowledge, and therefore have economic value to organizations. These skills, knowledge and experience represent capital because they enhance productivity. Motivations represent those psychological process that cause the arousal, direction and persistence of voluntary actions that are goal oriented. There are several motivation theories. Maslow's need theory was the development of the hierarchy of needs. He believed that there are at least five sets of goals which can be referred to as basic needs and are physiological, safety, love, esteem and self-actualization. Maslow stated that people, including employees at organizations, are motivated by the desire to achieve or maintain the various conditions upon which these basic satisfactions rest and by certain more intellectual desires.

One of the earliest researchers in the area of job redesign as it affected motivation was Frederick Herzberg. He and his associates began their initial work on factors affecting work motivation in the mid 1950's. Herzberg discovered that employees tended to describe satisfying experiences in terms of factors that were intrinsic to the content of the job itself. These factors were called motivators. Motivators (e.g. challenging work, recognition for one's achievement, responsibility, opportunity to do something meaningful, involvement in decision making, sense of importance to an organization) that give positive satisfaction, arising from intrinsic conditions of the job itself, such as recognition, achievement, or personal growth.

TABLE I.
HERZBERG-FACTORS

Motivators	Hygiene factors
satisfaction with the work	salary
challenge by the work	coworker relations
good team	safety
enough information	authority
	enough information
	work environment
	training, career
	spending leisure time
	solve problems, complaints

Conversely, dissatisfying experiences, called hygiene factors, largely resulted from extrinsic, non-job related factors, such as company policies, salary, coworker relations, and supervisory styles [3]. Hygiene factors (e.g.

status, job security, salary, fringe benefits, work conditions, good pay, paid insurance, vacations) that do not give positive satisfaction or lead to higher motivation, though dissatisfaction results from their absence. The term "hygiene" is used in the sense that these are maintenance factors. These are extrinsic to the work itself, and include aspects such as company policies, supervisory practices, or wages/salary.

II. THE RESEARCH

The aim of the research was to make a quality assessment system to Alba Regia Technical Faculty. In order to make that we create a questionnaire in which the motivations of the employee was estimated. The questionnaire was divided up to six parts.

In the first part there were personal data (birthdate, workplace-institute etc.). The second part was general, questions about the satisfaction with the workplace, the environment of the workplace etc. In the third part we assessed the motivations of the employees using Herzberg factors: intrinsic and extrinsic factors (motivators and hygiene factors). In the next part of the questionnaire we tried to estimate the talent of the employees. The aim of the next part was to measure the satisfaction with the teaching. The last part referred to the estimate the satisfaction with the assessment system.

24 colleagues filled the questionnaire, 60% was female, 40% male. 13% was leader, 87% was employee. The distribution of the birth can be seen in Table 2.

TABLE II.
DISTRIBUTION BY THE BIRTH

Birth date	Person	
1945-1964	11	46%
1965-1979	10	42%
1980-1994	3	12%

The colleagues ranked what is important for their on a workplace out of the next factors: salary, work environment, training, career, bonus (Fig. 1.)

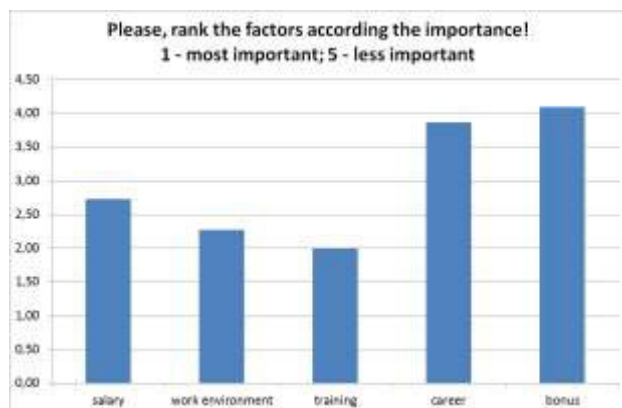


Figure 1. What is the most important for you in your workplace?

The survey was verified that for the colleagues is very important the training. The next factor was work environment: the most colleagues like to work alone (as teacher in general), and in a separate room. These things

are commonly characteristic the researchers. The third one was the salary, but it is true that as public servant it is not a motivator.

We got the same result in other analyses, in which we research what is the most important strong point in the workplace (Fig. 2.) The most important factor were the stability, the balance between the work and private life and the respect the colleagues. These factors are in essence the work environment. The less important factors were the outing and creature comforts.

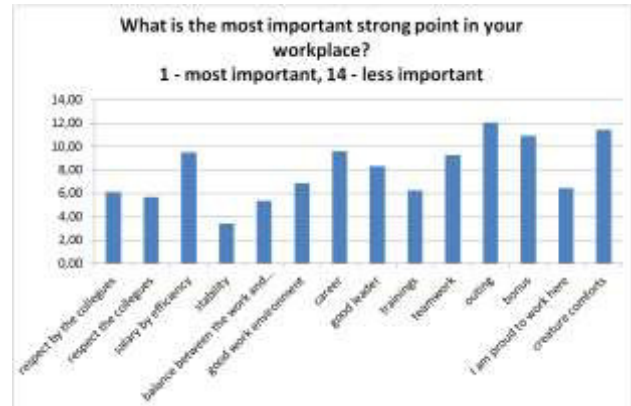


Figure 2. What is the most important strong point in your workplace?

To survey the motivation area among the colleagues were used the Herzberg factors. The questions were: How satisfied are you if..., and How dissatisfied are you if... (Table III, IV.)

TABLE III.
HYGIENE FACTORS

control	methods of the control
work specifications	technical specifications work environments
coworker relations	coworker relations relations with the leaders
salary and safety	work safety salary, bonus work-life balance
company policies	justice, legality

TABLE IV.
MOTIVATORS

achievement	deliverance
appreciation	direct feedback leader appreciation
work	work flow work diversity great importance work
responsibility	responsible work
career and growth	career

III. THE RESULTS

The first step was to analyses how dissatisfied the colleagues are if any factor is missing. In Table V. in the first column we can see the ranked statements (in Likert-scale 1-5). In the second column there are the average values of the statements, the next are the deviations (the average of the deviation is relatively high, 1,22); and in

the last column there are the sign of the Herzberg-factors (M-motivator, H-hygiene factor).

The survey verified us the hygiene factors, like safety, salary by efficiency, balance between the work and private life, the relationships with coworkers and leaders, work environment. The unfitting factor was achievement by the work, which can be due to the technologies and modern work logistic. Everybody work a lot and achieve in high level, so it is naturally if somebody good at work. The other factors were responsibility and great importance of work which are bound up with the achievement.

TABLE V.
DISSATISFACTION

How dissatisfied you if...	Average	Deviations	Factor
the work is uncertain	4,17	1,18	H
the salary is not by efficiency	4,08	1,11	H
not enough private life	4,04	1,06	H
not the best achievement	3,96	1,17	M
the coworker relationships are not good	3,88	1,33	H
not good relationship with the leader	3,71	1,51	H
not good work environment	3,63	0,99	H
you do not feel the importance of the work	3,63	1,32	M
not enough responsibility	3,50	1,35	M
not good company policies	3,42	1,15	H
you do not get the leader appreciation	3,38	1,18	M
no deliverance	3,38	1,28	M
not good control	3,25	1,33	H
the work is not diverse	3,25	1,20	M
no feedback	3,21	1,26	M
no career	2,96	0,93	M
the output of the work is hard to identify	2,96	1,34	M
lack of importance of the work	2,21	1,29	M
	3,48	1,22	

The lack of motivators does not cause dissatisfaction (below in Table V.) To control this statement we made another table (Table VI), in which we can find the answer to the question: How satisfied are you if....

We can wait that the above mentioned (Table V.) motivator factors will be in the Table VI. The first result is that the good achievement is the best important for the colleague. The other group of the motivators (career, feedback) has no high score both table, so these factors neither are motivators or hygiene factors, the employee do not calculate with these factors.

The method of the control and salary are hygiene factors. Motivators are responsibility and responsive work. Other motivators are leader appreciation and the good identified outputs.

The relationships with coworkers and leaders are motivators and hygiene factors also. That means, that

someone is missing, that can cause dissatisfaction, and if anything is very good, that can cause satisfaction. Also the balance with private life, the safety and work environment are motivators and hygiene factors also. Individuals spend a lot of time at work. It is safe to say that work may seem like their home away from home. Since they spend so much time with coworkers and management, they need a pleasant work environment.

TABLE VI.
SATISFACTION

How satisfied you if...	Average	Deviations	Factor
good achievement	4,67	0,75	M
good relationship with coworkers	4,54	0,76	H/M
good relationship with the leaders	4,46	1,00	H/M
safe work	4,42	1,04	H/M
the output can be easily identify	4,33	0,80	M
the work is diverse	4,25	0,78	M
good work environment	4,04	0,98	H/M
good company policies	4,04	0,93	H/M
enough private life	4,00	1,22	H/M
enough leader appreciation	3,96	1,10	M
good deliverance	3,96	1,21	M
responsible work	3,92	1,19	M
responsibility	3,92	1,08	M
salary by efficiency	3,92	1,22	H
good control methods	3,75	1,09	H
direct feedback	3,58	1,04	M
career	3,13	1,27	M
enough importance of the work	2,46	1,35	H
	3,96	1,04	

IV. CONCLUSION

Workplace satisfaction helps companies extract the best performance from their employees. According to Herzberg, hygiene factors are what cause dissatisfaction among employees in a workplace. In order to remove dissatisfaction in a work environment, these hygiene factors must be eliminated. There are several ways that this can be done but some of the most important ways to decrease dissatisfaction would be to pay reasonable wages, ensure employees job security, and to create a positive culture in the workplace. Eliminating dissatisfaction is only one half of the task of the two factor theory. The other half would be to increase satisfaction in the workplace. This can be done by improving on motivating factors. Motivation factors are needed to motivate an employee to higher performance.

The survey shows that the negative events have more influences to our satisfaction/dissatisfaction. That shows

the questionnaire and the tables above. Some suggestions are:

- Creating complete and natural work units where it is possible. An example would be allowing employees to create a whole unit or section instead of only allowing them to create part of it.
- Providing regular and continuous feedback on productivity and job performance directly to employees instead of through supervisors.
- Encouraging employees to take on new and challenging tasks and becoming experts at a task.
- Not receiving feedback on their work can be quite discouraging for most people. Effective feedback will help team members know where they are and how they can improve.
- Autonomy and control are necessary for people to feel satisfied with their work. In fact, psychologists have found that the less control people have over their jobs, the more stressful and unsatisfying they find it.

ACKNOWLEDGMENT

We acknowledge the financial support of this work by the Hungarian State and the European Community under

the TÁMOP-4.2.2.B-15/1/KONV-2015-0010 project entitled "Development of Alba Regia Technical Faculty's Scientific Workshops".

REFERENCES

- [1] P. Ramsden: Learning to teach in higher education. Routledge, 2003, 265 p. https://books.google.hu/books?hl=hu&lr=&id=As-4wAJz4C&oi=fnd&pg=PR1&dq=quality+assessment+method+higher+education&ots=XKXGRNkaK1&sig=JWWaxQngFHFg2phSbCcqyE5RrH0&redir_esc=y#v=onepage&q&f=false
- [2] A.W.Austin, A. L. Antonio: Assessment for excellence: The philosophy and practice of assessment and evaluation in higher education. Rowman and Littlefield, 2012, 367 p. https://books.google.hu/books?hl=hu&lr=&id=jEsRVtbukyMC&oi=fnd&pg=PR5&dq=benchmark+higher+education&ots=DJUbTSH33j&sig=vh5i5jNgxJkvt4KVk7imGyvqFKQ&redir_esc=y#v=onepage&q=benchmark%20higher%20education&f=false
- [3] S. Ramlall: A review of employee motivation theories and their implications for employee retention within organizations. Journal of American Academy of Business, Cambridge, Sep 2004, 5, ½. pp. 52. – 63. ftp://118.139.161.3/pub/moodledata/113/Ramlall_2004.pdf

The Modernization Opportunities of the Floodlight System of Prohászka Ottokár Memorial Church using Modern Technologies

Krisztián Kolozsvári*, Szilvia Kiss**

*Óbuda University Alba Regia Technical Faculty, Székesfehérvár, Hungary

**Óbuda University Alba Regia Technical Faculty, Székesfehérvár, Hungary

*kolozsvari.krisztian@hok.uni-obuda.hu

**kiss.szilvia@amk.uni-obuda.hu

Abstract- The Prohászka Ottokár memorial church has a 52-meter-high dome. When the church was consecrated in 1933, the dome was lit as so in 1938 in the year of jubilee. In the 1990's for several years the cross was lit on Sundays and holidays. However, nowadays the floodlight does not work in the church. The purpose of our research was to investigate the old floodlight systems. Our goal is to examine the possibilities of redesigning the old floodlight system using an environmental-friendly and cost-efficient LED technology. An advantage of the technology is that the system can provide a combination of colours during holidays.

Keywords: LED technology, Floodlight System

I. THE OLD LIGHTING METHODS

The Prohászka Ottokár memorial church was consecrated on 5th November 1933. During the ceremony the big and the small dome were equipped with floodlight as well as the gilded iron cross on the top of the small dome.

The church is located in the neighbourhood of the central railway station in Székesfehérvár. The church was lit for the first time in the night before the consecration ceremony because on 4th November 1933 on the occasion that the archbishop Jusztinián Serédi travelled through Székesfehérvár on the Roman pilgrimage train. The Fehérvár local newspaper reported on the ceremony on 4th November 1933. According to the newspaper article the floodlight was invented by the mayor of Székesfehérvár. He intended to illuminate the church and some other attractions such as statues, monuments and beautiful buildings in Székesfehérvár every Saturday and Sunday.

In 1938 the residents of the city celebrated the jubilee year of King St. Stephen. The church played a key role in the celebration because on 1st June 1938 the ashes of bishop Prohászka were transferred from the Long Cemetery to the Prohászka Ottokár memory church, which was built in memorial of bishop Prohászka. At the time the church was re-installed with floodlight. The dome was equipped with electric garlands in three wreaths.

A newspaper dated back to 19th May 1938 corresponds about the floodlight. According to the newspaper the residents of the city were so impressed that they wanted to

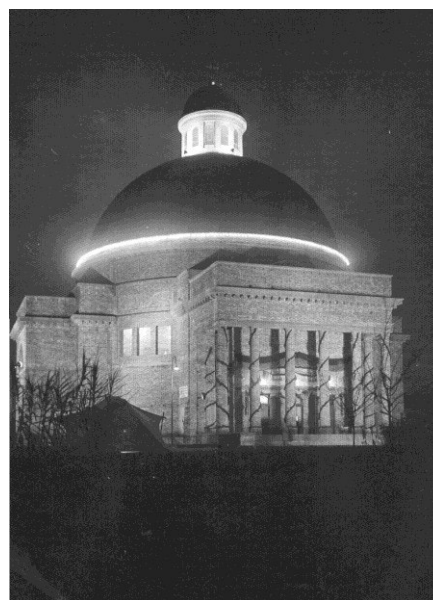


Figure 1. Prohászka Ottokár memorial church in 1938

see the floodlight every Saturday and Sunday thereafter. Therefore it is assumed the vision in 1933's cannot be realized.

Another news item we found about the floodlight references to 2nd April 1939 when the anniversary of bishop Prohászka's death was celebrated.

The last sentence of the article mentions lights and the hundreds of light bulbs.

Five years later, on 3rd April 1944 when the city celebrated another anniversary of bishop Prohászka's death, the local newspaper mentioned that the church could not be lit up due to the fact that the country was under war.

On 19th September 1944 the army sieged the city of Székesfehérvár and the church was hit by several bombs. The copper cladding and the wooden dome structure were badly damaged. The dome caught fire and blazed for 24 hours and wood parts were burnt. In 1945 the church was hit by about 40 grenades. The dome damaged heavily. The renovation began only in 1946 with restoring the big dome in 1948. However, the overall renovation of the dome was

completed in the summer of 1961. From this period we could not find any written notes about the floodlight system.

We revealed that in the 1980's the parish attempted to



Figure 2. The dome burned for 24 hours

partially install a floodlight system. Nevertheless, only the cross was lit by two spotlights on the top of the church.

In this period the parish priest called Joseph Garda issued a regulation. According to this regulation the floodlight was supposed to be operated every Saturday, Sunday and during religious holidays.

The floodlight can be switched on and off manually inside the church.

The church was under reconstruction in 2002. In this time the electrical system was completely renovated. The church was equipped with new electric wires and luminaries. The cross on the top of the big dome was equipped with two spotlights. From this time on only the cross has been possible to be illuminated.

Other designing and modernization efforts aiming the external floodlight were cancelled and the lighting in the cross continues to be operated only manually.

Over the years the spotlights went wrong, so the church has not been floodlighted outside for more than ten years. [1-6]

II. PLANS FOR A NEW ILLUMINATION SYSTEM

We did research with the aim of finding old documents and photos in the church and by means of this information we decided to reconstruct the old floodlight system. Our goal is to equip the big dome's bottom, the small dome's bottom, and, of course, the cross with lighting. A long time ago the people used simple electric garlands, hundreds of bulbs and a manual on off switching system for the implementation.

In our opinion it is necessary to automatize the floodlight system in an energy saving and environmental friendly way. We recommend using the LED lighting systems because they are energy saving, environmental friendly solutions and more resistant than traditional light bulbs.

In the events of religious and state holidays we intend to change the colours of lights in line with the nature of the holidays.

III. ADVANTAGES AND DISADVANTAGES OF LED LIGHTING

The LED lighting has a number of advantages.

The biggest advantage of LED lighting is the energy-saving operation, which is especially important in our case, because we want to shed light on a large surface area applying a cost-effective operation.

A. Advantages of led lighting:

- **Light output:** In case of traditional light bulbs efficiency is very low, because much of the energy is converted to heat, approximately 8-14lm / W, while the efficiency of LED lighting is very good about 50-80lm / W.
- **Endurance:** The lifetime of conventional bulbs is around 1000 hours, whereas the LED bulb's life is approximately between 80 and 100 thousand hours of operation.
- **Low maintenance costs:** The long lifetime of LED lighting systems requires a one-time investment, so the realization of the floodlight system in the memorial church is beneficial, because the system will be installed hard to reach places on the façade and at the bottom of the dome.
- A further advantage is that the system can be operated in small places, the light of Led does not vibrate and it is not bad for the eyes.
- **Light controllability:** Led lights up where needed and will not cause light pollution.
- If it is necessary the brightness and colours of Led can be controlled.
- **Switchable without delay:** It is not necessary to wait for warm-up; Led can be lit immediately with appropriate light intensity.
- **On / Off - switches:** The high number of switches does not shorten the life span.
- **Low Voltage operation:** The operating voltage of LED strip lighting system is approximately 12-24V; therefore this system offers a high-level safety operation.
- **Light spectrum of the LED:** The led's light spectrum is narrow, therefore, led does not contain UV or infrared ranges, as a consequence, the light of led cannot damage the material and colours of objects.

B. Disadvantages of led lighting:

- **High cost of installation,** although it pays off in the long term considering the maintenance and the energy costs.

At the end of our research, we wanted to test our concept regarding the facade of the church in an area where we could place a 5-m long LED strip.

We wanted to operate the partial floodlighting for a couple of nights and get feedback from the parish and the residents in the neighbourhood, thus promoting further implementation of the system, possibly supported by EU funds as well.

Although we could not get all the permits that we needed for the test before the article submission deadline, the realization of our vision is in progress. In terms of

testing, we have created a simple microcontroller-operated circuit. With this circuit we can switch the red and white led strips on and off in twilight automatically.

IV. THE LED STRIP'S PROPERTIES THAT WE USED IN OUR PROJECT

The used LED strips are as follows: - 5050 LED strip of 30 LEDs / m = 420 lumens light output, 7.2 Watt performance. The high-performance LED strips are produced from 5050 mm SMD LED. The name comes from the size, which is 5.0X5.0 mm.

Comparing the 5050 led to the other led strips a larger physical size; larger performance, larger consumption and larger heat emission can be seen. This is due to the fact that inside the 5050 mm 3 pieces SMD LED light emitting chip is located. Theoretically the brightness and the consumption are three times bigger than in case of 3528 led strips.

The led strips that we planned for the church floodlight have IP68 level of protection. These led strips are equipped with a two-layer protective coating.

The first layer of the resin casing is IP65, and the second layer is 1 mm square tubes, which are closed airtight and watertight.

The double insulation provides protection even if the outer casing is damaged. These LEDs do not include adhesive tape layers, thus fixing the fastening of the clips is to be solved.

In case of the church floodlight system it is especially important, since the installation of led strips has to be solved properly because the installation procedure to the high dome is very expensive



Figure 3. LED strip with mounting clip

In the third figure you can see the LED strip along the retaining clip.

V. THE CIRCUITS FOR TESTING

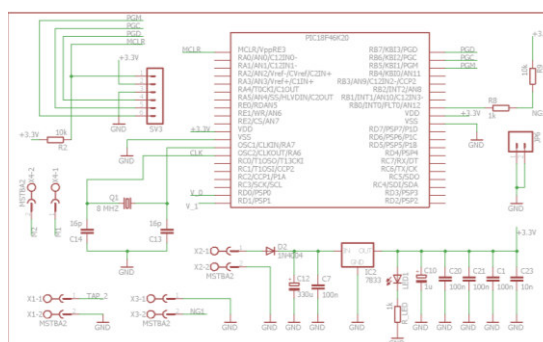


Figure 4. The implementation of testing circuit-microcontroller

In the fourth picture the implementation of testing circuit's microcontroller parts can be seen.

The function of a part of the circuit is the following. By pushing a button we can choose the red or the white led strips to be switched on. The V_0 sign permits the red and the V_1 sign permits the white led strips operation.

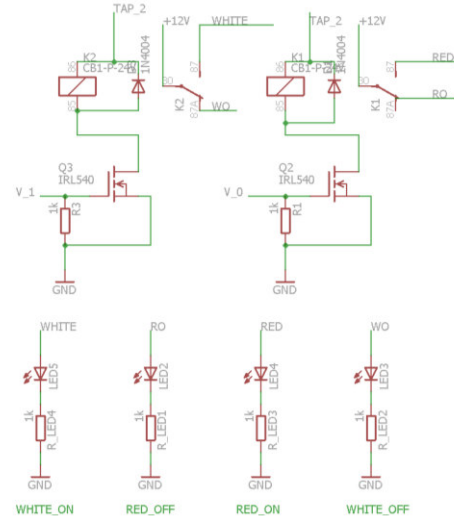


Figure 5. The enable signal help to drive a relay

The fifth picture you can see the enable signal is drive to a relay. The circuit has been placed with 4 LED indicator where you can check which colour mode is selected.

After the relay we separately built two dark detector circuits to the led strips. With the help of this circuit we can set up the level of light where white led strips and red led strips must be activated. The solution can be seen in the picture 6.

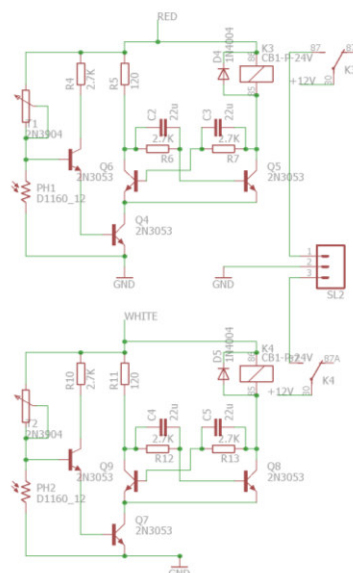


Figure 6. Dusk sensing circuit

VI. SUMMARY

In summary, we believe that our topic is verified by the fact that the residents who have been living in the neighbourhood of the old church for many years would like to see the illuminated domes and the cross on the top of the building.

Our aim was to design an innovative way to effectuate the floodlight system that has long-life and low-maintenance costs, which might be a suitable alternative for the implementation and can be tested involving the environment to correct this.

We hope that within a few years the parish can find adequate resources to install the led strips floodlight system in the church and our idea will be realized. [7-11]

ACKNOWLEDGMENT

The research work is supported by Óbuda University and the Prohászka Ottokár Memorial Church.

REFERENCES

- [1] Kultusz és emlékezet, 1927-1938. (Kiadó: Székesfehérvár Megyei Jogú Város Levéltára, Székesfehérvári Püspöki és Székeskáptalani Levéltár, 2013)
- [2] Szemelvények templomunk történetéből VII. (Kiadó: Vasútvidéki Prohászka Egyházközség, 2001. szeptember 7.)
- [3] Szemelvények templomunk történetéből VIII. (Kiadó: Vasútvidéki Prohászka Egyházközség, 2001. szeptember 23.)
- [4] Szemelvények templomunk történetéből IX. (Kiadó: Vasútvidéki Prohászka Egyházközség, 2001. szeptember 30.)
- [5] Szemelvények templomunk történetéből X. (Kiadó: Vasútvidéki Prohászka Egyházközség, 2001. október 21.)
- [6] Szemelvények templomunk történetéből XI. (Kiadó: Vasútvidéki Prohászka Egyházközség, 2001. november 18.)
- [7] Beágyazott mikroprocesszoros rendszerek tervezése és fejlesztése (Összeállította: Polyák Béla, KKVMF, 1990)
- [8] Mikrovezérlők alkalmazás technikája: PIC mikrovezérlők (Kónya László, ChipCAD, 2000)
- [9] Közvilágítási kézikönyv (MEE Világítástechnikai Társaság Magyar Világítástechnikáért Alapítvány, szerkesztő: Kosztolicz István)
- [10] <http://energiapedia.hu/>
- [11] <http://www.anrodiszlec.hu/>

The Application of Fuzzy Logic in the Field of Land Consolidation

János Katona

Óbuda University Alba Regia Technical Faculty, Székesfehérvár, Hungary
katona.janos@amk.uni-obuda.hu

Abstract— In Hungary land consolidation lacks a plenty of preconditions. The information based decision-making model belongs to these preconditions. This model would be able to optimize the present land structure based on objective aspects. The basis of the consolidation can be the Goldencrown value of the property registry. Since this value is out-of-date, correction factors should be set out. Geoinformatics can contribute to spatial data managing. What is more, fuzzy logic can also provide solution for the description and summation of factors. The present paper gives an example for the description and mapping visualization of modifying factors using sigmoid function on a sample area.

I. THE NECESSITY OF LAND CONSOLIDATION

The present fragmented land structure originated in the distribution of compensation and joint ownership property happened in the 1990s. The privatization of field land was applied to the $\frac{3}{4}$ of the arable of the country. 5.6 million hectares were shared among 2.6 million private individuals so that the average land size per capita was about 2 hectares. These land areas were often distributed in parcels and/ or as undivided joint ownership. With this land reform that group in the society received land that did not want to deal with agriculture. Nowadays, in Hungary, 64% of the owners hire out their land property. Therefore the property and the use of land have been significantly diverged. For that reason, in Hungary, land lease – contrary to West-European countries – is not the sign of the spread of competitive farms, but of the unambiguous consequence of privatization [1].

Since the change of the political system some kind of land unification has been observable in Hungary. Land size is a significant part of competitiveness. So the 8.1 hectares average land size [Fig. 1], which has come into existence due to land unification, is still little as compared to other European models [2].

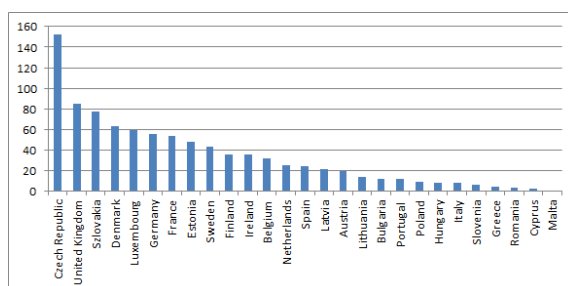


Figure 1. Average land size in the European Union

The Act No. CXXII of 2013 concerning agricultural and forestry land trade has aimed to solve the problem of fragmented land structure. The law supports individual and family farmers as against national and international land speculators. The importance of the regulation is unquestionable, however, the spectacular results can be observed only in the long run. The government has also regulated (374/2014. regulation) the elimination of the undivided joint ownership started up in the course of land privatization. 950.000 hectares land area and 300.000 applicants are involved in the project. Although this act can solve the problem of the undivided joint ownership, it still increases the fragmentation of land structure.

A complex land consolidation would be necessary both for the approximation of land use and land ownership and land unification. According to a comprehensive paper [3] published in 2004, the institutionalized land consolidation could be put into practice. On the one hand, the conditions of such a consolidation are given: there is a land registry system, skilled labour and technical knowledge. On the other hand, there are conditions that can come into existence such as legal regulation, budget, socialization and an appropriate IT solution that can be generally applied in the course of land consolidation. The aim of the present paper is to introduce an IT solution that can be used to determine the exchange value in the course of land consolidation. During land consolidation, it is not necessary to express the exchange value in sale value; the use of modified GoldCrown, that would involve the most essential aspect modifying the value, could be sufficient.

II. FUZZY LOGIC AND LAND CONSOLIDATION

Fuzzy logic is significant because it can describe the imprecise and imperceptible information in a quantitative way. On the one hand, the fuzzy “set” can be seen as a classical mathematical set which means that there exist numbers that are either the elements of the set or not.

$$\mu(x) = 1, \text{ if } x \text{ is the element of } X$$

$$\mu(x) = 0, \text{ if } x \text{ is not the element of } X$$

On the other hand, the fuzzy function assigns a value to each and every element, which can describe the extent to which the element belongs to the set. The value of the fuzzy function can be 0-1.

$$F = \{ (x, \mu_F(x)) : x \in X \}$$

$$\mu_F(x) \rightarrow [0,1]$$

To lower the capacity of calculation, linear functions have spread. But it is important to know that fuzzy functions can occur in any higher form. Nowadays fuzzy

logic is widely used such as in the field of automotive sector, medical technology, factory farms or entertaining technology [4].

As for land consolidation, fuzzy logic can contribute to estimation of the value in exchange. Due to the lack of land consolidation law, we need to take into account of 54/1997 FM regulation to define the factors affecting valuation of property. According to the regulation, the corrective factors are these: shape, size, location, reach, road conditions, relief and gradient conditions, drainage, landmarks interfering cultivation, esthetical impression, the likelihood of frost, ice and game damage, irrigation, inclination for cultivation, demographical proportions, nutritional farming, agrochemical intervention, environmental pollution and long-term environmental impairment, natural protection, melioration.

The above-mentioned factors always need to be taken into account in the course of valuation of property. However, on the behalf of the social acceptance of land consolidation, it is not practical to take into consideration many factors that could affect the GoldCrown value on a large scale.

Table I. contains the most important corrective factors with their intervals and the definition of evaluation.

TABLE I.
CORRECTIVE FACTORS

Corrective factors	Interval		The basic of valuation
	lower	upper	
Shape, area, size	-20	20	Area, perimeter
Location	-20	20	Distance from the clear
Reach and road condition	-20	20	Categories of connected roads
Relief and gradient conditions	-20	0	Categories of slopes
Melioration and drainage	-20	0	Likelihood of flood and polder
Conditions of irrigation	0	20	Distance from the irrigational channel
Environmental protection	-15	0	On the basis of land use limitation

The fuzzy functions can be made on the basis of the intervals of corrective factors. Since, the functions can have only values 0-1, the neutral value should be 0.5 in any cases. The sigmoid function has been selected for fitting to the accurate series. This kind of function is useful because it is possible to give both the neutral value and the value of the slope of the sinus curve.

$$mSig(x,a,c) = 1/(1 + \exp(-5/a * (x - c)))$$

, where:

a=the growth of slope

c=neutral value 0,5

The shape factor, that can affect valuation of property, needs some explanation. Since the value of the property is influenced by how it can be cultivated, it is important to determine the extent of correction in an objective way. It is well-known that among plane figures the ratio of the area and the circumference gives the biggest number. This value is influenced by the squared variable derived from two dimensions. Therefore the square can be eliminated

and we can get a value that is appropriate for comparison. The value of the property is affected by the growth of a continuous area that is why the above-mentioned value needs to be corrected. The area size needs to be raised to 0.1 power to make the shape and the area size almost equal.

$$\text{Shape factor} = (\sqrt{\text{area/perimeter}}) * \text{area}^{0,1}$$

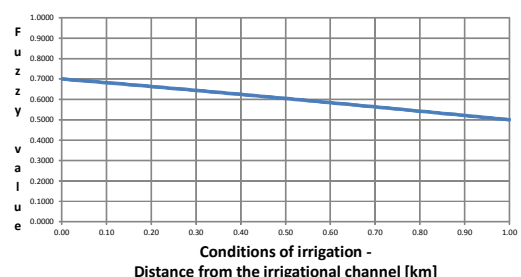
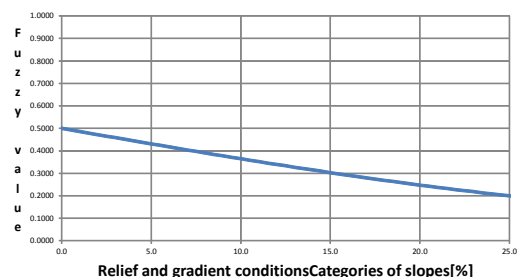
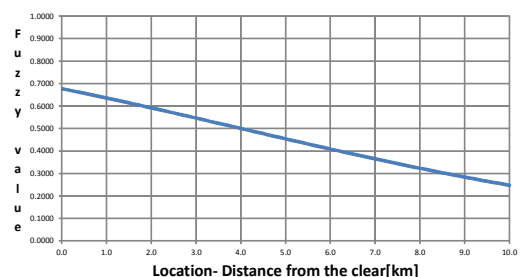
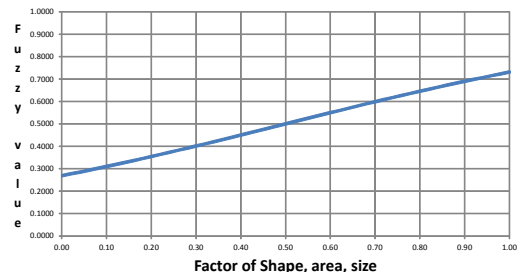


Figure 2. Fuzzy functions

Fig.2 shows the functions of different factors affecting the property. The functions of some factors seem to be linear but to the extreme converging –not seen- sinus curve parts are significant, too, because multiplying by a number greater than 1 or smaller than 0 can be problematical in the course of summation of corrective factors.

III. ANALYSIS OF THE SAMPLE AREA

The formation of the value mapping of the field land expands in Mesterszállás. The village is the sample area of

the Institute of Geodesy, Cartography and Remote Sensing. The Institute has provided the land registry data of the village. The land registry contained 1578 parcels, out of which 566 cultivated parcels have been selected for the analysis. The average size of the parcels is 5.8 hectares, the quality classes are among 2-7. Fig.3 shows the original GoldCrown values.



Figure 3. GoldCrown values on the sample area

Table II. summarizes the results of the analysis containing 6226 records.

TABLE II.
STATISTICS OF FACTORS AFFECTING VALUATION OF PROPERTY

Corrective factors	Min	Max	Average	Range	Scatter
Shape, area, size	0.32	0.70	0.50	0.38	0.0770
Location	0.32	0.68	0.56	0.36	0.0832
Reach and road condition	0.50	0.60	0.51	0.10	0.0287
Relief and gradient conditions	0.50	0.50	0.50	0.00	0.0000
Melioration and drainage	0.30	0.40	0.30	0.10	0.0094
Conditions of irrigation	0.56	0.70	0.68	0.14	0.0309
Environmental protection	0.45	0.50	0.50	0.05	0.0051
Sum	0.45	1.62	0.94	1.17	0.2070

The above mentioned data show how the different kinds of factors have influenced the valuation of the analysed

parcels. There are significant differences among the corrective factors; however, they are not surprising in view of the sample area. The smallest extent corrective factor can be found in case of relief. The reason for this is that the difference between the highest and the shortest point of the village does not reach 10 metres; therefore there are no measurable differences in slope. The largest correction can be observed in case of “Shape, area and size”. This conclusion is in accord with the result coming from a representative survey among land owners.

According to this survey, the above mentioned factor influences the unit value eminently after the GoldCrown value. After the summation of the corrections, the discrepancy between the smallest and the biggest correction is 117%. The smallest correction (0.45) was given to a parcel (its topographical number 0188/7) which has worse shape and size factors than the average; it is the farthest from the inland, it can be approached on dirt road, the likelihood of flood damage is high, it has average irrigational conditions and it does not stand under environmental protection. The highest corrective factor (1.62) was given to a parcel (its topographical number 0312/4) which has better shape and size factors than the average: it is near the inland; it can be approached on pitched road, the likelihood of flood damage is high, it has average irrigational conditions and it does not stand under environmental protection. Fig.5 shows the value of totalized corrections – the values are shaped after normal distribution.

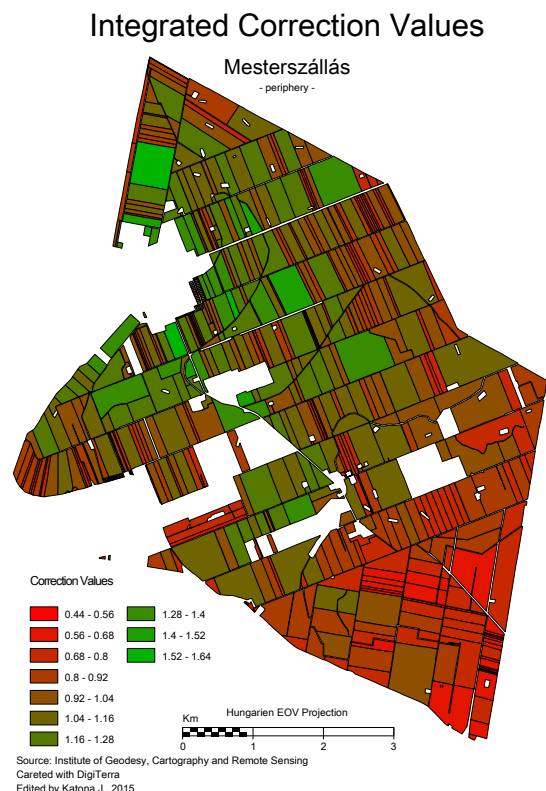


Figure 4. Integrated Correction values on the sample area

The analysis principally leans on cadastre mapping substance. Two corrective factors (shape, area, size, location, geographical position) have been elicited from this data resource. What is more, this substance has

provided the mapping basis of other corrections. Free available digital resources (such as Interior Ministry General Directorate of Water Management - Floodmap, Rural Development Ministry – Nature Conservation Information System) have also contributed to the determination of other factors.

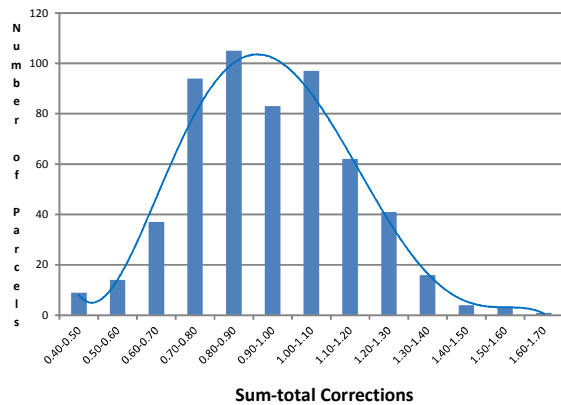


Figure 5. Statistics of Sum-total Correction

IV. SUMMARY

To carry out successful land consolidation, social support is necessary. The development of the IT solution mentioned in the present paper serves this aim. The statistical weighting of certain factors can be easily modified by the change of the slope values of Fuzzy functions. Therefore it can be geared to the demand of the community that is affected by the consolidation. Naturally, an institutionalized land consolidation has got several conditions and it also needs trial projects that analyse the technical realization. The following part of the research dealing with land consolidation aims to use the results of the evaluation method that applies fuzzy logic.

ACKNOWLEDGMENT

The described work was carried out as part of the TÁMOP-4.2.2.B-15/1/KONV-2015-0010 project as a work of Development of science education workshop at Alba Regia Technical Faculty in the framework of the New Széchenyi Plan. The realization of this project is supported by the European Union, cofinanced by the European Social Fund.

REFERENCES

- [1] J. Alvincz - R. Schmidt: A birtokrendezés főbb kérdései Magyarországon, különös tekintettel a földcserére, *Geodézia és Kartográfia* vol. 10. 2008, pp 26-32.
- [2] J. Ripka: Birtokszerkezet, Nemzeti Birtokrendezési Stratégia és a földügyi szakigazgatás, *Geodézia és Kartográfia* vol. 9. 2005, pp. 23-30.
- [3] L. Dorgai et al.: A Magyarországi Birtokstruktúra, A Birtokrendezési Stratégia Megalapozása, *Agrárgazdasági tanulmányok* vol. 6., Agrárgazdasági Kutató Intézet, Budapest, 2004, p.199.
- [4] T. Varga: A fuzzy logika alkalmazási lehetőségei a minőségtervezésben, *Debreceni közlemények* vol. 1. 2010, pp. 43-51

A New Way in Business Intelligence

Why and how Oracle's In-Memory database could have the ability to change the Business Intelligence based decision support systems?

Author: Patrik Szpisják, Co-author: Dr. Levente Rádai

Óbuda University, Alba Regia Technical Faculty, Székesfehérvár, Hungary

Óbuda University, Alba Regia Technical Faculty, Székesfehérvár, Hungary

E-mail: szpisjak.patrik@hok.uni-obuda.hu, radai.levente@amk.uni-obuda.hu

Abstract

Nowadays, almost every company use various IT technics to improve their decision making systems. The most common is to use analytic technics like OLAP. After the analytic process, we have to represent the data, that's where BI comes in place. While OLAP is a single high performance analytic tool, BI is a description of all the processes, tools and methodologies about how to access the information. It contains reporting, OLAP, dashboards, score carding, DWH and other products to extract the information. Currently, BI is called the most modern way to maintain the success of a company. But is this enough?

I. INTRODUCTION

Today, it's the era of Big Data. Since the 1980s, the world's technological per-capita capacity to store information has roughly doubled every 40 months. At 2012, 2.5 Exabyte of data was created each day. That is a huge amount of information. We have to work with more and more data each day, therefore we need to change the technology in order to keep up with the expanding needs. An easy way is to change the hardware, not the amount but the technology in order to improve the performance. This was where the In memory databases came in mind [1].

II. IN MEMORY DATABASES

In memory databases (IMDBs) are generally designed to be fast but until now, DRAMs were so expensive that it was unaffordable to move the whole database into memory for almost any company. Today, terabytes of data can be stored in memory in a cost of some thousand dollars. With the data in memory, access times has been highly reduced and performance has been greatly improved. Besides the advantages there are some complexities with this architecture. Recoverability is critical in all databases. If the machine goes down, you must not to lose any of the data.

Therefore you have to write the data to the disk as well as placing it in memory. It is cost effective but doing I/O operations on disk is slow and you can hardly write fast enough to keep up with the data that is being putted into memory. The other problem with this technique is that you risk the data consistency. Since the I/O operations on disk is slow, after some time these operations will fall behind and you may lose some data when the in memory database shuts down. Mirroring could be a solution. Duplicating it on two different machines keeps the database consistent. The safety level is also higher but if the two machines go down at the same time, you lose all the data. Using the combination of these two is complex and costly. There are other ways to handle the problem like Cloud Computing but each one of them only offers partial solution [2].

III. ORACLE IN-MEMORY DATABASE

Oracle's In-Memory Database is part of the latest released database, Oracle Database 12C. It is not a database itself, more like an extension. Oracle In-Memory Database allows you to store the data in disk and in memory at the same time without becoming inconsistent. When an insert commits, the row will be inserted traditionally to the disk. Meanwhile, if the parameter has been set on the appropriate table where the data is being inserted, another thread will pump the data into the memory. This thread only wakes up every 2 minutes in order to keep the data consistent [3].

While the data on the disk is stored in traditionally row format, In-memory it is stored in column format in order to compress it. The compression level depends on the data compress algorithm we set. Oracle offer 3 major choices. The basic compress is no compress where the data costs as much space as on the disk. Another one is for query performance which can be separated into 2 parts: for query

low and for query high. The “for query low” option makes some compression while preferring optimizing for queries. “For query high” do less compression in order to maximize the queries performance. Depending on the data itself, these two options can make a 2-10 time compression ratio.

The third one is for capacity reduction. As before, you can choose among two sub categories: for capacity low and for capacity high. The “for capacity low” makes much compression but still cares with the queries performance. “For capacity high” does the most compression of all, almost totally neglecting performance. The ratio is really high here, depending on the data, it could even reach 100-1000 times compression ratio. Since the ratio is so high, in order to use queries on the appropriate object, first the database has to decompress it. Another mentionable thing about compression is that you can not only apply it on tables but also on subset of tables. This allows you to use for capacity high on rarely needed columns and for query high on often used columns.

Oracle built in an adviser called compression adviser. This feature estimates the compression ratio of the given object by virtually compressing it.

Column format also allows the CPU to use vector processing. This format is specifically designed to maximize the column entries number that can be evaluated in a single CPU instruction. This enables to scan billions of rows per second per core instead of the capacity of the buffer cache, which can scan only millions of rows per second per core [3].

You can set yourself which object you would like to put into memory. It is possible to set this option on tablespaces, tables, columns and even parts of the columns, example on partitions too. Therefore it is highly customizable for the queries. Setting the in memory size allows you to decide what amount of data would you like to store in memory. In case the memory is full, population to the memory stops and when some space frees up in memory, it will be pumped in.

Oracle offers to use priority levels on in memory objects. The possible 5 levels are none, low, medium, high and critical. The critical data is on the highest level and the priority decreases as we go backwards. The goal is to always keep the critical data in memory, so when we “startup” the instance, these data will be populated first followed by high priorities and so on until we populate all or as many data as we can. It is important that this option cannot be used on subset of tables [4].

Recoverability did not change. Since you have all the data on disk, after shutting down and starting up, everything will

recover on your disk and the appropriate ones will be populated back into the memory.

Besides reporting and ad-hoc queries, there is also an improvement in the Online Transaction Processes, called OLTPs. To increase the performance, the designer had to avoid the excessive use of indexes. The cost of maintaining indexes are high, but since we can highly customise how and what we want to store in memory, it is possible to drop the analytic indexes on the tables that is stored on the disk. Of course, it is necessary to keep the primary and foreign keys.

IV. BUSINESS INTELLIGENCE NOWADAYS

In order to effectively grant fast access to data, we use various technics and technologies under the name of Business Intelligence. Using OLTP technics to pump the generated data to the database is fast and effective. This data is inserted on the disk in row format therefore it can be extracted slowly because the slow I/O operations of the disk. On the other hand, it is not well structured for applying analytic queries. The best way is to use an Extract Transform Load (ETL) process. During this phase, the data is being extracted from the original database, being transformed to the required format – which can be a column store format too by pivoting – and then being loaded into another database which is technically a data warehouse (DWH). A typical ETL based DWH uses multiple layers: staging, data integration and access layers usually called data marts. The staging layer stores all the data extracted from the disparate data source systems. The data integration layer transforms the data from the staging layers into data sets. These data sets are usually divided by operational systems like sales, financial and marketing. Finally this integrated data is arranged into hierarchical groups often called dimensions and into facts and aggregated facts in a typical form of star form. This is the layer from which we extract the data with various reporting tools [5].

Since each layer could be a database itself, building a BI system could cost much money by buying database licenses, support, various tools, hiring ETL specialists and database administrators. Besides money, it also costs much time. Developing a well optimized and structured BI system for reporting could take up to two years or even more. Another disadvantage is if the data warehouse shuts down, no data can be accessed which may cause huge delays in important decisions. Looking at performance, by transforming the data into the required format roughly decreases the generating time of reports, dashboards, etc. There is a real possibility of having to wait for minutes until it finally generates. This is caused by the slow I/O operations of the disk [6].

V. ORACLE IN-MEMORY DATABASE AS A REAL-TIME ANALYTIC TOOL

The importance of a BI system and the need of various reports, dashboards and other analytics is unquestionable. In order to keep up with the market trend changes, it is a must to improve the decision making systems by using the new IT technologies. Using Oracle In-Memory feature on the data warehouse could roughly improve the performance or in some cases it can even eliminate the need for a data warehouse which was required to achieve fast analytic execution. Yet the real power is reflected upon using it on the live transactional data while also accelerating OLTP [7]. The ETL process which loads the data warehouse runs at nights when the OLTP system is less busy in order not to affect the transactional system's performance. Therefore, the analytic queries does not represent the current state. With this feature activated, real-time ad-hoc analytics is possible, enabling leaders to get on-the-fly answers. Real-time analytics is the answer when the timeliness of the insights is important or when you need a very large influx of live data. Applying this technology on an organisation, it could detect the errors instantly, after changing business strategy changes can be noticed immediately, fraud can be detected, keeping up with customer trends becomes possible, can improve service and gives better sales insights which could lead to additional revenue [8].

In my research, I would like to measure the performance difference between a single Oracle Database and an Oracle In-Memory Database in two fields: BI and real-time analytics. On the BI field, I will use two different queries on a DWH, both with in memory disabled and with in-memory enabled. In real-time analytics, I will use two queries again to measure the performance difference on the unstructured data and the number of columns selected, as it is suggested in [9]. The compression algorithm was set to for query low during all tests.

VI. RESEARCH, TESTING

The first test runs on a traditional BI system which is based on a data warehouse. In order to show the difference, I used two types of analytic queries. The first query selects three columns, one from the fact table and two from the joined two dimension tables. The fact table has 13 columns, the dimension tables have 3-6 columns. The selected columns are: department name, sales price and country name. From these, sales price will be aggregated.

A. First query:

```
select sum(a.price), b.department_name, c.country_name
from sales a, dep b, country c where
a.department_id=b.department_id and
```

```
a.country_id=c.country_id group by b.department_name,
c.country_name order by c.country_name ASC;
```

On the first test, dep table contains 27 rows, country table contains 4 rows and the sales table contains about 15 million rows.

Results:

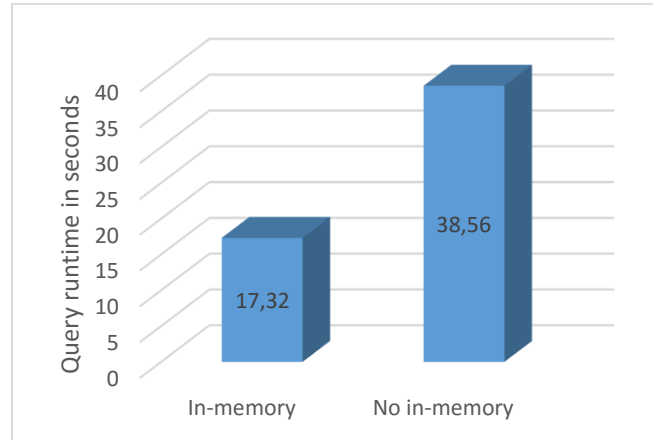


Figure 1. Using the first query: sales table has about 15 million rows

Executing the query a getting the results from the traditional database took 38.56 seconds. After enabling in-memory feature and populating the data into memory, the runtime of the same query was only 15.32 seconds. The runtime when in-memory enables is two times less when disabled.

Doubling the number of rows, the result looks like this:

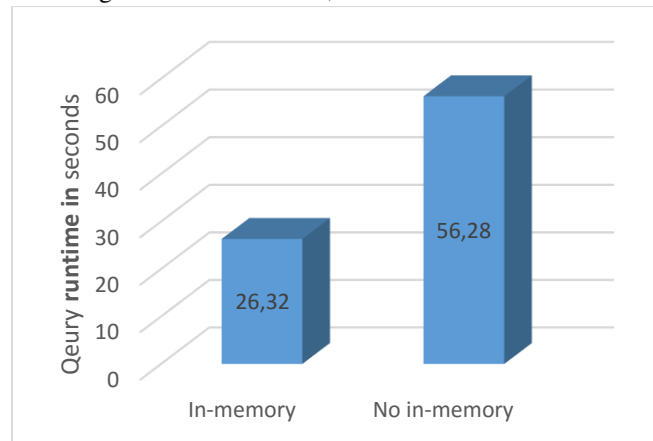


Figure 2. Using the first query: sales table has nearly 35 million rows

In this case the sales table contains nearly 35 million rows. As we see, it took more time to get the results than before. The number of rows increased the runtime, but did not double it. On the next performance test, I will increase the number of joins from two to five and the number of column from 3 to 6.

The second query selects six columns, one from the fact table and five from the joined five dimensional tables. The columns are: department name, country name, material type, payment type, sales price and region name. As before, sales price will be aggregated.

B. Second query:

```
select sum(a.price), b.department_name, c.country_name,
d.material_type, e.payment_type, f.region_name from sales
a, dep b, country c, mat d, paym e, region f where
a.department_id=b.department_id and
a.country_id=c.country_id and a.material_id=d.material_id
and a.paym_id=e.paym_id and a.region_id=f.region_id
group by b.department_name, c.country_name,
d.material_type, e.payment_type, f.region_name order by
c.country_name ASC;
```

In this test, dep table has 27 rows, country table has 4 rows, mat table has 14 rows, paym table has 6 rows, region table has 9 rows and the sales table has about 15 million rows.

Results:

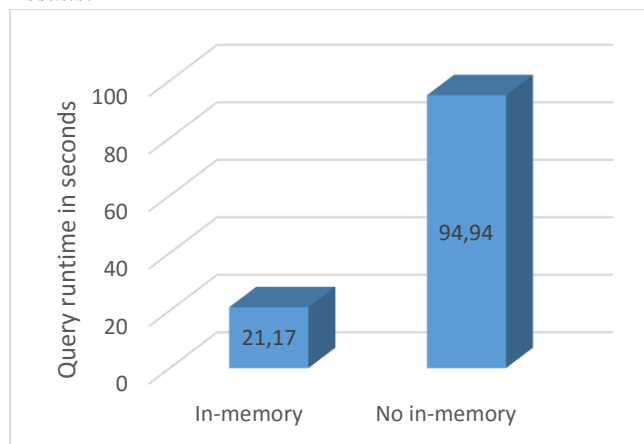


Figure 3. Using the second query: sales table has about 15 million rows

In this test, we can see that by joining more dimensional tables to the fact table, the query runtime significantly increased. This time, the query ran about 4-5 times faster when the in-memory feature was enabled and the data was populated into the memory.

We have seen before, that doubling the rows affected the performance, however joining more dimension tables increased the runtime way more.

Double the number of rows and see the results again:

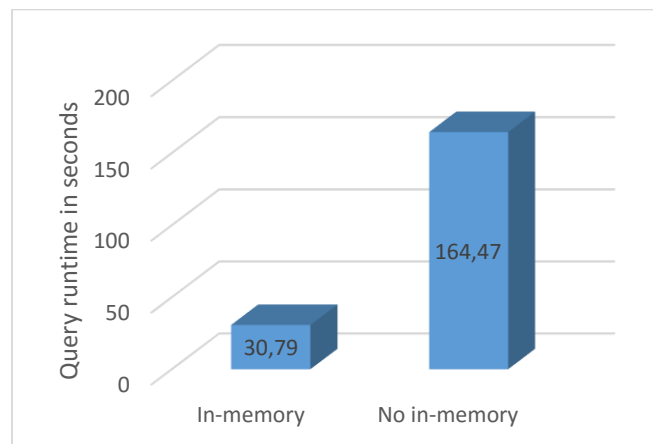


Figure 4. Using the second query: sales table has nearly 35 million rows

Generally speaking, using more joins significantly decreases the performance, while using the query on a bigger table decreases performance firmly less. On average, 2 dimensional queries are 2x faster while 5 dimensional queries are 6x faster with the in-memory feature enabled. The more complex the query, selecting the data from in-memory becomes more effective.

We can clearly say, that using data warehouses with in-memory option roughly increases performance and reduces time needed to generate various reports. In the era of Big Data, it is a powerful tool which helps in day by day decisions. It can speed up the preparing of business processes which could lead to market benefits.

The second test is represents a real-time analytics. The most important difference is that in the OLTP systems, the data is not structured well and is not transformed to the required format. Therefore, using analytic queries can cause overloads, can shut down systems which could lead to market loss. Using a powerful analytic tool which does not affect the OLTP performance is a key to real-time analytics. The first test showed how in-memory benefits after we insert more data and do more joins. In this test, I will use two queries which will concentrate on the performance issues caused by the unstructured data and the number of columns selected. The query selects six columns, one from the fact table and five from dimension tables. The fact table has nearly 70 columns, the dimension tables have 3-11 columns. The selected columns are the same as on the other test: department name, country name, material type, payment type, sales price and region name. Like before, sales price is aggregated again. There are some differences on the table and column names, but the query is the same as in the second.

C. Third query:

```
select sum(a.price), b.dep_name, c.country_name,
d.mat_type, e.paym_type, f.reg_name from sales a, depart b,
country c, mater d, paym e, region f where
a.dep_id=b.dep_id and a.country_id=c.country_id and
a.mat_id=d.mat_id and a.paym_id=e.paym_id and
a.region_id=f.region_id group by b.dep_name,
c.country_name, d.mat_type, e.paym_type, f.reg_name
order by c.country_name ASC;
```

In this test, depart table has 27 rows, country table has 4 rows, mater table has 14 rows, paym table has 6 rows, region table has 9 rows and the sales table has more than 90 million rows, nearly 100 million. I chose this large number of rows, because this will surely show the difference.

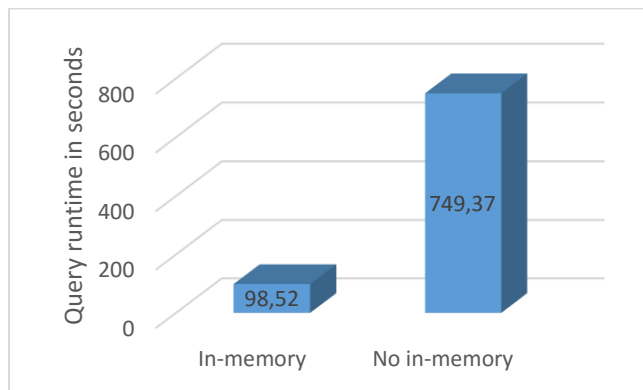


Figure 5: Using the third query on unstructured data: sales table has nearly 100 million rows, 6 columns selected

The difference is huge. By the advantage of vector processing and in-memory column store architecture, using in-memory on unstructured data improves the performance even more than before.

During the second test, I will use the same tables but I will select 18 columns instead of six.

D. Fourth query:

```
select sum(a.price), a.date, a.load_id, a.ref_number
b.dep_name, b.dep_id, c.country_name, c.contry_id,
d.mat_type, d.mat_id, d.mat_unit, d.seq_num, e.paym_type,
e.paym_div, e.load_id, f.reg_name, f.code, f.id from sales a,
depart b, country c, mater d, paym e, region f where
a.dep_id=b.dep_id and a.country_id=c.country_id and
a.mat_id=d.mat_id and a.paym_id=e.paym_id and
a.region_id=f.region_id group by a.date, a.load_id,
a.ref_number b.dep_name, b.dep_id, c.country_name,
c.contry_id, d.mat_type, d.mat_id, d.mat_unit, d.seq_num,
e.paym_type, e.paym_div, e.load_id, f.reg_name, f.code,
f.id order by c.country_name ASC;
```

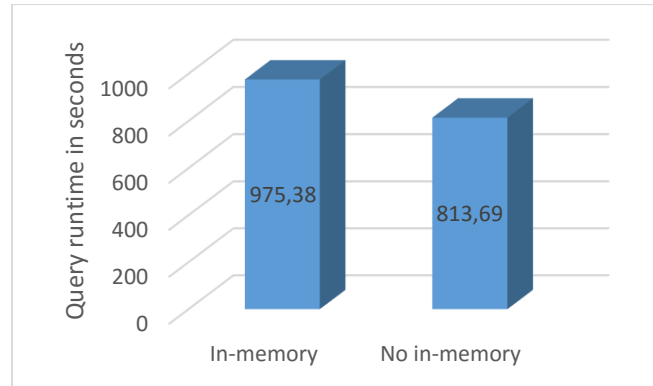


Figure 6: Using the fourth query on unstructured data: sales table has nearly 100 million rows, 18 columns selected

This is where in-memory column store architecture is not effective. Upon selecting a larger number of columns, the performance highly drops. This is caused by the in-memory column store architecture which was designed to select only a few columns.

Generally speaking, in order to benefit from Oracle In-Memory Database we have to keep in mind that selecting large number of columns and returning large number of rows can cause huge performance loss. Furthermore, using complex predicates and joining multiple large tables has bad impact on performance. Therefore it is reasonable to look into the details of tests and create a performance difference map based on the number of rows and columns.

VII. SUMMARY

Oracle In-Memory Database has the ability to change the regular BI concept. Being able to store the required data in-memory for real-time ad-hoc analysing while also accelerating OLTP on disk is powerful. Performance benefits can be also seen at aggregation, at filtering and at BI type reporting which typically summarises a lot of data. Nevertheless, in most cases, if we need the data instantly from huge tables it does not enough, we still need aggregated tables. Therefore, if the organization creates and stores a huge amount of data using this feature instead of a data warehouse is not recommended. The real power shows upon using it on a single database for real-time ad-hoc querying. Depending on the settings and on the data, queries could be executed a hundred times faster or even more. At last, having the chance to use one product for multiple purposes offers great support possibility and prevents problems which can appear when integrating products from separate organizations.

REFERENCES

- [1] István Orosz, Tamás Orosz, Company level Big Data Management, In: Anikó Szakál (szerk.), 9th IEEE International Symposium on Applied Computational Intelligence and Informatics ; SACI 2014. Conference place, date: Timisoara, Románia, 2014.05.15-2014.05.17. (IEEE), Timisoara: IEEE Hungary Section, 2014. pp. 209-303.
- [2] A Selmecei, T Orosz, Usage of SOA and BPM changes the roles and the way of thinking in development, In: Szakál A (szerk.), 2012 IEEE 10th Jubilee International Symposium on Intelligent Systems and Informatics, SISY 2012, Subotica, 2012, September, 20-22. Konferencia helye, ideje: Subotica, Szerbia, 2012.09.20-2012.09.22. Piscataway: IEEE, 2012. pp. 265-271.
- [3] Oracle Database In-Memory, Oracle Database 12c, Oracle White Paper, Oracle Corporation, July 2015, URL: <http://www.oracle.com/technetwork/database/in-memory/overview/twp-oracle-database-in-memory-2245633.html>, p: 29, Downloaded: 15th September, 2015.
- [4] Oracle Database In-Memory, Powering the Real-Time Enterprise, Oracle Data Sheet, Oracle Database 12c, Oracle Corporation, URL: <http://www.oracle.com/technetwork/database/options/database-in-memory-ds-2210927.pdf>, p: 9, Downloaded: 15th September, 2015.
- [5] Tamás Orosz, István Orosz, Optimization of Enterprise Business Processes with Big Data Management, BULETINUL STIINTIFIC AL UNIVERSITATII POLITEHNICA DIN TIMISOARA ROMANIA SERIA AUTOMATICA SI CALCULATORAE / SCIENTIFIC BULLETIN OF POLITECHNICA UNIVERSITY OF TIMISOARA TRANSACTIONS ON AUTOMATIC CONTROL AND COMPUTER SCIENCE 59:(2) pp. 105-110. (2014)
- [6] A. Selmecei, I. Orosz, Gy. Györök, T. Orosz, Key Performance Indicators Used in ERP Performance Measurement Applications, In: Szakál A (szerk.), 2012 IEEE 10th Jubilee International Symposium on Intelligent Systems and Informatics, SISY 2012, Subotica, 2012, September, 20-22. Konferencia helye, ideje: Subotica, Szerbia, 2012.09.20-2012.09.22. Piscataway: IEEE, 2012. pp. 43-48.
- [7] L. Ellison and J. Loaiza, Oracle Database In-Memory - Powering the Real-Time Enterprise, Oracle YouTube Channel, URL: <https://www.youtube.com/watch?v=CjkPUg1m6PA>, Published: 13th June, 2014.
- [8] C. Millsap, K. Osborne and T. Poder, Oracle Database InMemory in Action, URL: https://www.youtube.com/watch?v=Ze3umRJ_HBM, Published: 12th November, 2014.
- [9] M. Rittman, In-Memory Analytics Speeds Business Intelligence, Oracle Magazine, January/February 2015, URL: <http://www.oracle.com/technetwork/issue-archive/2015/15-jan/o15ba-2398995.html>, Downloaded: 15th September, 2015.

On a problem concerning rainbow Ramsey theory

Borbély József, Csala-Takács Éva
Óbudai Egyetem, Székesfehérvár,
Hungary

borbely.jozsef@amk.uni-obuda.hu

csala.takacs@amk.uni-obuda.hu

Abstract:

In this paper we will present a problem from the branch of the so called rainbow Ramsey theory. The problem was posed by András Gyárfás for the annual Schweitzer Miklós Competition in 2009. The solution of the problem is based on the argumentation of Borbély (who submitted a solution at the competition). This solution was refined by the co-author Csala-Takács.

Introduction

Ramsey theory is a part of mathematics that handles problems where one seeks after ordered structures in some great arbitrarily chosen structures. In a more restricted sense the problems of Ramsey theory have abstract or even geometric graphs as their subjects. In this context a Ramsey type theorem includes an allegation, in which one guarantees the existence of some “ordered” subgraphs in every according graph great enough. In order to clarify the nature of problems treated in Ramsey theory, we hereby draft Ramsey’s classical coloring theorem concerning abstract graphs.

Ramsey used this theorem (in another equivalent form, see [2]) to achieve results in the formal logic.

Some Ramsey-type theorems

Here we will present some classical Ramsey-type theorems in order to demonstrate the nature of these problems. First we present Ramsey’s classical result (see [2]):

Theorem 1 (Ramsey, 1930) *For every positive integers k and s there exists a positive integer $R(k, s)$, such that the following is true: if one colors the edges of a complete graph G with at least $R(k, s)$ vertices with red and blue, then there is a complete subgraph of G with k vertices having all its edges colored by red, or there is a complete subgraph of G with s vertices having all its edges colored by blue.*

Later this theorem was rediscovered and had a many other applications. In a wider sense one can speak about the so called Ramsey principle (Ramsey type results in other structures).

The theorem can be generalized for more colors in an analogous way.

Theorem 2 (Ramsey, generalized form) For every positive integers $s_1, s_2, \dots, s_k$ there exists a positive integer

$R(s_1, s_2, \dots, s_k)$ so that the following is true: if one colors the edges of a complete graph G having at least $R(s_1, s_2, \dots, s_k)$ vertices with k colors, then there is a complete subgraph of G with s_i vertices having all its edges colored by color i for some i .

There exists even an infinite version of Ramsey's theorem.

Theorem 3 (Ramsey, infinite form) Let X be some countably infinite set and colour the elements of $X^{(n)}$ (the subsets of X of size n) in c different colours. Then there exists some infinite subset M of X such that the size n subsets of M all have the same colour.

One of them most substantial theorems of Ramsey theory is Van der Waerden's theorem.

Theorems of the rainbow Ramsey theory have the goal to guarantee great structures colored by distinct colors. This is a developing branch of mathematics with a wide range of problems (see [1]).

Our main theorem

We present the problem of Gyárfás that was posed in 2009 at the Schweitzer competition.

Theorem 4: If $c_1, c_2, \dots, c_m$ are colorings of the edges of a complete graph with 17 vertices (labeled by the numbers $1, 2, 3, \dots, 17$) by 105 fixed colors in such a way, that for every 15-element subset H of the vertices there is a j so that in the coloring c_j all of the 105 colors occur in the coloring of the edges of H , then $34 \leq m$. Moreover the desired conditions can be fulfilled for $m=34$.

Proof:

A complete graph with 17 vertices has exactly $\binom{17}{2} = 136$ edges. First we will prove that every complete

graph with 17 vertices, whose edges are colored by 105 colors, contains at most 4 complete subgraphs H with 15 vertices whose edges are colored by the 105 colors.

Note that such a complete subgraph H has also the property that every edge of H has a different color (a rainbow property). Let us call such complete subgraphs rainbow-subgraphs. Now, with this terminology, we have to prove that every complete graph with 17 vertices, whose edges are colored by 105 colors, contains at most 4 rainbow-subgraphs.

Let us take an arbitrary coloring of the edges of the complete graph with 17 vertices by 105 colors.

We will have two different cases.

Case A,

If there is a 16-element subset of the vertices where we can find at least two rainbow-subgraphs, then there is a set F of the vertices with 14 elements and two vertices x and y such that the graph spanned by the vertices of $F \cup \{x\}$ and $F \cup \{y\}$ are rainbow-subgraphs.

Note that in this case the color classes represented between the sets $\{x\}$ and F must be the same as the ones represented between $\{y\}$ and F . Therefore the edges whose endpoints are x or y represent at most 17 color classes (namely there are 31 edges with the endpoints x or y , but there are 14 color classes that are used two times, and $31 - 14 = 17$).

This means that x and y cannot be the vertices of a rainbow-subgraph at the same time. Therefore there are exactly two rainbow-subgraphs spanned by the vertices of the set $F \cup \{x\} \cup \{y\}$. Let z be the seventeenth vertex of the complete graph with 17 vertices. If there are another rainbow-subgraphs, then their vertex set must contain z . Assume that there is such a rainbow-subgraph. In this case $\{z\} \cup F$ determines a rainbow-subgraph, or there is an element u in F so that $(F \setminus u) \cup \{z\} \cup \{x\}$ or $(F \setminus u) \cup \{z\} \cup \{y\}$ determines a rainbow-subgraph.

If $\{z\} \cup F$ determines a rainbow-subgraph, then the color classes represented by the edges running between F and z are the same as the ones running between F and x and the ones running between F and y . According to our observations it is clear that no two elements of the set

$\{x, y, z\}$ can be vertices of a rainbow-subgraph at the same time, therefore there are totally at most 3 rainbow-subgraphs.

If $(F \setminus u) \cup \{z\} \cup \{y\}$ spans a rainbow-subgraph, then the color classes represented by the edges running between z and $(H \setminus u) \cup \{y\}$ are the same as the ones running between u and $(H \setminus u) \cup \{y\}$. Using our argumentations before one can easily verify that the vertices u and z cannot be the vertices of a rainbow-subgraph at the same time.

The same is true for the vertices x and y . Every rainbow-subgraph contains at least two vertices of the set $\{x, y, z, u\}$, but not more than two, because the elements of the sets $\{x, y\}$ and $\{u, v\}$ exclude each other. Therefore there are at most 4 rainbow-subgraphs in Case A,

Case B,

If every subset of the vertices with 16 elements defines at most one rainbow-subgraph, then let us take a 16-element subset of the vertices where there are 15 vertices determining a rainbow-subgraph.

Let us take another rainbow-subgraph, too (if there is not any more one, then we are ready). In this case there is a subset K of the set of the 17 vertices with $|K|=13$ and vertices r, s, t, v so that the graphs defined by the sets $\{r\} \cup \{s\} \cup K$ and $\{t\} \cup \{v\} \cup K$ are rainbow-subgraphs.

Then the color classes represented by the edges running between the vertex sets $\{r\} \cup \{s\}$ and K are the same as the ones running between $\{t\} \cup \{v\}$ and K (26 colors, totally). The edges that have r, s, t , or v as endpoints, can represent at most 32 color classes, because

$$\binom{4}{2} + 26 = 32.$$

This means that the vertices r, s, t, v cannot be the elements of a rainbow-subgraph at the same time. Therefore there are at most 4 rainbow-subgraphs as we have stated.

Using our result we have the inequality $4m \geq \binom{17}{15}$,
which is equivalent to $34 \leq m$.

Now we will prove that there are colorings $c_1, \dots, c_{34}$ fulfilling the requirements of the theorem.

Let L be a 13-element subset of the vertices and let a, b, c, d be the other vertices of the graph.

Let us prepare a rainbow-subgraph on the vertex set $\{a\} \cup \{b\} \cup L$. Let the color classes between c and L be the same as the ones between a and L , and let the color classes between d and L be the same as the ones between b and L . Let the color of the edge between b and c the same as the one between b and a . Let the color of the edge between a and d the same as the one between a and b .

In this case the vertex sets $\{a\} \cup \{b\} \cup L$, $\{a\} \cup \{d\} \cup L$, $\{c\} \cup \{b\} \cup L$ determine rainbow-subgraphs. If we color the edge connecting c and d by the color connecting a and b , then $\{c\} \cup \{d\} \cup L$ is a rainbow subgraph, too. We will give such colorings in all of the 34 cases. It is enough to choose the ordered quadruples (a, b, c, d) in an appropriate way.

Remember that the vertices of the graph are labeled by the numbers 1, 2, ..., 17. We will order to every rainbow-subgraph a two element subset in such a way, that we

identify the rainbow-subgraph with the two vertices that are not the vertices of it.

With this notation and our construction, for every quadruplet (i, j, k, l) we are able to establish a coloring so that the pairs $\{k, l\}$, $\{j, k\}$, $\{i, l\}$, and $\{i, j\}$ represent a rainbow-subgraph. In 17 cases we will make the coloring so, that $(i, j, k, l) = (x, x+2, x+1, x+4)$, where x runs through the elements of the residue classes modulo 17.

$$k-l = -3 \pmod{17},$$

$$l-k = 3 \pmod{17}$$

$$j-k = 1 \pmod{17},$$

$$k-j = 1 \pmod{17}$$

$$i-l = -4 \pmod{17},$$

$$l-i = 4 \pmod{17}$$

$$i-j = -2 \pmod{17},$$

$$j-i = 2 \pmod{17}$$

Therefore we get all the pairs of indices that have the difference (depending on the order) 1, 2, 3, 4, -1, -2, -3, -4 modulo 17.

We will make other 17 colorings, using this pattern. In this case we will have $(i, j, k, l) = (x, x+6, x+1, x+8)$, where x runs through the elements of the residue classes modulo 17. In this case we have

$$k-l = -7 \pmod{17},$$

$$l-k = 7 \pmod{17}$$

$$j-k = 5 \pmod{17},$$

$$k-j = -5 \pmod{17}$$

$$i-l = -8 \pmod{17},$$

$$l-i = 8 \pmod{17}$$

$$i-j = -6 \pmod{17},$$

$$j-i = 6 \pmod{17}$$

Therefore we get all the pairs of indices that have the difference (depending on the order) 5, 6, 7, 8, -5, -6, -7, -8 modulo 17.

This completes the proof of our theorem.

With the same method one can prove the following more general result:

Theorem 5: *If $c_1, c_2, \dots, c_m$ are colorings of the edges of a complete graph with $n=8k+1$ vertices (labeled by the numbers $1, 2, 3, \dots, 8k+1$) by $\binom{n-2}{2}$ fixed colors in such a way, that for every $(n-2)$ -element subset H of the vertices there is a j so that in the coloring c_j all of the colors occur in the coloring of the edges of H , then*

$$\frac{1}{4} \cdot \binom{n}{2} \leq m.$$

References

- [1] R. Graham, B. Rothschild, J. H. Spencer, Ramsey Theory (2nd ed.), New York: John Wiley and Sons, 1990
- [2] V. Jungic, J. Nešetřil, R. Radoicic, Rainbow Ramsey theory, Integers, 2005, A09, 13 pages
- [3] F. P. Ramsey, On a problem of formal logic, Proceedings of the London Mathematical Society, 1930, 264–286.

Rough Set Based Efficiency Improvement of Simulation Performance Prediction

László Muka and István Derka

Széchenyi István University, Department of Telecommunications, Győr, Hungary

e-mails: muka@sze.hu, steve@sze.hu

Abstract— Parallel and distributed simulation method (PADS) has been showing growing importance in the discrete event simulation (DES) analysis of large-scale, complex network systems. To realize a good performance with the PADS method, it is necessary to predict the behavior of the simulation model in the parallel and distributed execution environment. This paper describes a Rough Set Theory based method of simulation performance prediction efficiency improvement using Systems Performance Criteria (efficacy, efficiency and effectiveness). The new improved performance prediction method – Algorithm of Improved Simulation Performance Prediction (AISPP) – is presented as addition to the traditional Coupling Factor Method (CFM) prediction approach.

Keywords— Parallel discrete event simulation, simulation performance prediction, Rough Set Theory, Systems Performance Criteria, Coupling Factor Method

I. INTRODUCTION AND MOTIVATION

Over the last few years, together with the increase of the need for the discrete event simulation (DES) analysis of large-scale, complex systems and networks requiring high computing capacity, a lot of efforts have been made in research of parallel and distributed discrete event simulation modelling and execution methods [1-4], since the parallel and distributed execution simulation turned out to be an appropriate approach to improve simulation runtime performance [5, 6].

The parallel and distributed simulation can be defined in different ways. According to a simpler but broader view, parallel and distributed simulation is any simulation in which more than one processor is used [3]. According to another accepted definition, the Parallel Discrete Event Simulation (PDES) [6] and the Parallel and Distributed Simulation (PADS) [3] are differentiated on the feature, whether the single simulation model is executed on a set of *tightly* coupled processors (e.g. a shared memory multiprocessor) or on a set of distributed *loosely* coupled processors (e.g. PCs interconnected by LAN or WAN). According to a simple view, parallel and distributed simulation (PADS) is any simulation in which more than one processor is used [3]. In the present paper, PADS is defined as the execution of a single discrete event simulation model on high performance computing platforms (like clusters of homogeneous and heterogeneous computers), and on emerging platform environment (WEB, grid and cloud).

Despite of its beneficial effect on the simulation runtime performance, the PADS method is not in the

everyday use in the simulation community, because *development* of simulation models having high runtime performance features in a parallel and distributed execution environment remained a hard task even today [7]. Simulation *performance prediction* methods and tools can help to realize higher performance with the PADS model development by providing *preliminary* knowledge about the likely *behavior* of the model [7-11]. The PADS performance prediction methods should support the performance analysis throughout the development and evaluation process [7].

Summary of the overview above can be as follows: uncertainty of data necessary for the simulation performance modelling and prediction is present in every phase of the simulation process, and both the uncertainty and cost of obtaining data is increasing with the distance (measured in simulation cycle-steps) of prediction from the simulation execution.

The *motivation* of the authors to make the research presented in the paper was the lack of methods which can manage together the accuracy of performance predictions – the basic requirement for performance prediction – and the cost of data observations necessary for parallel and/or distributed simulation performance prediction.

For the new method presented in the paper the *Rough Set Theory (RST) method* and the systems approach of performance – *Systems Performance Criteria (SPC) efficacy, efficiency and effectiveness* have been selected to handle together the accuracy and the cost of predictions in a train-and-test analysis.

In the present paper, the authors made the following major *contributions*:

The RST model of simulation performance prediction has been defined allowing train-and-test analysis of performance predictions

A set of *operations and measures* are given for the use of *efficacy (E1)*, *efficiency (E2)* and *effectiveness (E3)* to measure for a single and for a series of predictions in order to embed SPC into RST train-and-test process.

The algorithm of the new *Algorithm of Improved Simulation Performance Prediction (AISPP)* has been described. Including the *Coupling Factor Method (CFM)* of the parallel and/or distributed simulation performance prediction in the AISPP process a feedback is provided to support the model identification and refinement stage.

The rest of the paper is organized as follows. The RST, SPC and CFM approaches are introduced in Section 2. In Section 3, the new prediction improvement method is

formulated: first definitions of E1, E2 and E3 are given for calculations and attribute and rule dropping then the new AISPP process is described. Section 4 concludes the work.

II. INGREDIENTS OF THE PREDICTION EFFICIENCY IMPROVEMENT APPROACH

A. Rough Sets

The RST (Rough Set Theory) is a mathematical framework particularly suitable for modelling and analysis of information systems with imprecise relations, with uncertain, vague data [25-29].

A *rough set information system* with embedded knowledge consists of two sets: the set of *objects* called the *universe* and the set of *attributes*.

More formally, $I = (U, A, f, V)$ denotes an *information system* of RST, where set U is the universe, A is the set of attributes. Sets U and A are finite nonempty sets where $(U = \{x_1, x_2, x_3, \dots, x_{|U|}\})$ and $A = \{a_1, a_2, a_3, \dots, a_{|A|}\}$. The attributes define a *transformation function* $f: U \rightarrow V$ for U where the set V is the set of values of A ($V = V_{a_1} \cup V_{a_2} \cup V_{a_3} \cup \dots \cup V_{a_{|A|}}$). The set V_{a_i} – named also the *domain* of a_i – contains the collection of values of a_i and $V_{a_i} = \{v_{1a_i}, v_{2a_i}, v_{3a_i}, \dots, v_{|V_{a_i}|a_i}\}$ where $|V_{a_i}|$ is the size of the domain of a_i .

Discretization is the operation of mapping the primary values and ranges of all attributes to *selected* (possibly optimized) sets of discrete values: $f_V: V' \rightarrow V$. The set V' stands for the values of a before discretization.

The *B-indiscernibility relation* $IND(B)$ for a set of attributes $B \subseteq A$ is defined in the following way:

$$IND(B) = \{(x_i, x_j) \in U^2 \mid \forall (a \in B) (a(x_i) = a(x_j))\}$$

If $(x_i, x_j) \in IND(B)$, then the objects x_i and x_j are indiscernible from each other in B and the *equivalence classes* $[x]_{IND(B)}$ of $IND(B)$ are formed by the objects indiscernible in B .

Rough sets are defined by their lower approximation and upper approximation sets. The set $B^*(X)$ and the set $B_*(X)$ is the *B-lower* and *B-upper approximation* of the set X and defined as follows:

$$B_*(X) = \bigcup_{x \in U} \{x \mid [x]_{IND(B)} \subseteq X\}$$

$$B^*(X) = \bigcup_{x \in U} \{x \mid [x]_{IND(B)} \cap X \neq \emptyset\}$$

The set $BN_B(X)$ defined by the equation $BN_B(X) = B^*(X) \setminus B_*(X)$ is the *B-boundary region* of X . If X is a crisp set then, $X = B_*(X)$ thus $BN_B(X) = \emptyset$ which means the boundary region is empty.

A *reduct* R_B is the minimal subset of attributes B that allows the same classification of objects of U as the set of attributes B . This feature of a reduct may be described by indiscernibility function as follows:

$$IND_{\forall x(x \in U)}(R_B) = IND_{\forall x(x \in U)}(B), B \subseteq A.$$

In general, the information system may take the form of $I = (U, A = C \cup D, f, V)$ which is a *decision information system (DIS)*. The set $C = \{c_1, c_2, c_3, \dots, c_{|C|}\}$

denotes the set of condition attributes and D is the set of decision attributes $D = \{d_1, d_2, d_3, \dots, d_{|D|}\}$. The information function $f: U \rightarrow V$ may be expressed by information functions $f_C: C \rightarrow V_C$ and $f_D: D \rightarrow V_D$, where $V = V_C \cup V_D$ ($V_C = V_{c_1} \cup V_{c_2} \cup V_{c_3} \cup \dots \cup V_{c_{|C|}}$ and $V_D = V_{d_1} \cup V_{d_2} \cup V_{d_3} \cup \dots \cup V_{d_{|D|}}$) and

$$V_C = \bigcup_{i=1}^{|C|} V_{c_i} \text{ where } V_{c_i} = \{v_{1c_i}, v_{2c_i}, v_{3c_i}, \dots, v_{|V_{c_i}|c_i}\}$$

$$\text{and } V_D = \bigcup_{i=1}^{|D|} V_{d_i} \cdot V_{d_i} = \{v_{1d_i}, v_{2d_i}, v_{3d_i}, \dots, v_{|V_{d_i}|d_i}\}.$$

In a *decision table* $I = (U, A = C \cup \{d\}, f, V)$ based on a DIS, d denotes the *distinguished decision attribute*.

Furthermore, a decision information system having the form of $I = (U, C \cup D, f_V, f, V', V)$ denotes a DIS with *discretization information function* $f_V: V' \rightarrow V$ (which is identical to information functions $f_{V'_C}: V'_C \rightarrow V_C$ and $f_{V'_D}: V'_D \rightarrow V_D$).

The *classification* may also be described by a *set decision rules* ($S = \{s_1, s_2, s_3, \dots, s_{|S|}\}$) in the form of implication ($s_i = (\varphi_i \Rightarrow \kappa_i), s_i \in S$), where φ_i and κ_i are *logical expressions* of the condition and decision attributes respectively. The formulas φ_i and κ_i may also be quoted as *LHS (Left Hand Side)* and *RHS (Right Hand Side)* part of the rule. A decision rule s_i may be evaluated using its *coverage* and *accuracy* ratios according to formulas

$$Coverage_U(s_i) = \frac{|Match_U(s_i)|}{|U|}$$

$$Accuracy_U(s_i) = \frac{|Supp_U(s_i)|}{|Match_U(s_i)|}$$

where $Match_U(s_i)$ is the number of objects in U the attribute values of which satisfy φ_i (*matching* with the *LHS* part of s_i), and $Supp_U(s_i)$ denotes the number of objects in decision table the attribute values of which satisfy both φ_i and κ_i (*matching both* with the *LHS* and *RHS* parts of s_i).

B. Systems Performance Criteria

The three *Systems Performance Criteria (SPC)* are the efficacy (E1), efficiency (E2) and effectiveness (E3) [11,30,31]. The SPC are in a hierarchy-like relationship with each other. On the longer term, the performance of a system is checked by the effectiveness criterion, the efficacy criterion shows whether the performance is suitable at all, and the efficiency criterion characterizes the relation of the required output and the resources used to produce the output.

C. The Coupling Factor Method

Based on some theoretical considerations about the connectedness of PADS model segments, paper [14] describes a practical simulation performance prediction approach the *Coupling Factor Method (CFM)*. The method – using results that have been got in sequential simulation runs – predicts the *parallelization potential* of simulation models (for models with conservative null message-based algorithm) and formulates requirement on how this potential can be exploited.

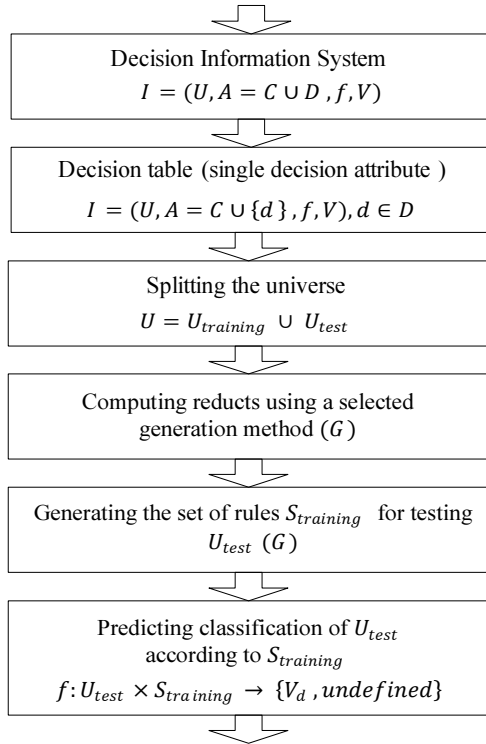


Figure 1. Traditional RST (TRSTA) algorithm for train-and-test analysis and learning

The principle of CFM may be summarized in an inequity:

$$L * Ev \gg \tau * P$$

where L is the lookahead value characterizing the model (simsec), E is the event density generated by the model (ev/simsec), τ is the latency of messages between logical process (LPs) of the model (sec), and P is the event processing computation hardware performance (ev/sec). In this practical approach, parameters L and E characterise the model itself, parameters τ and P describe the execution environment. According to the method, the coupling factor λ is calculated according to the formula $\lambda = L * E / \tau * P$. The high value of the coupling factor λ shows the good potential for simulation model parallelization. The method involves only four parameters for the performance prediction calculations. These parameters can be measured in simple *sequential simulation runs*. For a separate process, the λ_N parallelization potential of a process is only a part of the whole potential:

$$\lambda_N = \frac{\lambda}{N_{LP}} = \frac{1}{N_{LP}} * \frac{L * E}{\tau * P}$$

where N_{LP} the number of the LPs [15].

The method has been validated by a series of simulation experiments for *homogeneous* and *heterogeneous* clusters of computers [15-17].

Examples for the telecommunications networks and cloud computing systems are introduced in [16,32].

III. THE NEW PREDICTION EFFICIENCY IMPROVEMENT METHOD

The new method is built around the *Traditional RST Analysis Algorithm (TRSTA)* and the *E1, E2 and E3 Systems Performance Criteria*. (Steps of TRSTA train-and-test analysis process are shown in Figure 1.)

For the simulation *performance prediction efficiency improvement*, SPC can be identified as follows: to achieve the necessary prediction quality (E1 criterion), to realize it with an acceptable cost (E2 criterion) and to produce it in a stable manner on long run (E3 criterion).

A. Prediction performance calculations and dropping criteria

In case of a simulation performance analysis *objects* of the universe are *computer simulation experiments*. In a decision table $I = (U, A = C \cup \{d\}, f, V)$ attributes A describes the explanatory (independent) and dependent variables, decision rules $S = \{s_1, s_2, \dots, s_{|S|}\}$, $s_i = \varphi_i \Rightarrow \kappa_i$ describes the classification of the object of the experiments.

In the following, functions are defined allowing to calculate the E1, E2 and E3 criteria and supporting the increase of performance by dropping attributes and rules.

DEFINITION 1. CLASSIFICATION PREDICTION FUNCTION

$$f_{(U_{test}, S_{training})}: U_{test} \times S_{training} \rightarrow \{V_d, undefined\}$$

$$f_{([c](x_l), \varphi(s_k))} = d(x_l)_{predicted} = \begin{cases} \langle d(s_k) \rangle & \left(Match_{U_{test}}(s_k) = 1, (s_k \in S_{training}), \right. \\ & \left. \langle d(s_k) \rangle = v_d \in V_d \right. \\ & \left. undefined \mid otherwise \right. \end{cases}$$

DEFINITION 2. MATCHING PREDICTION OPERATOR

$$\bigvee_{l=1}^{|U_{test}|} (x_l) (x_l \in U_{test}) ([\langle c_{l,1} \rangle, \langle c_{l,2} \rangle \dots \langle c_{l,|C|} \rangle])$$

$$\xrightarrow{\bigvee_{k=1}^{|S_{training}|} Match(s_k)} [d(x_l)_{predicted}]$$

DEFINITION 3. PREDICTION CORRECTNESS

Prediction correctness $p(x_l)$ is

$$\forall (x_l) (x_l \in U_{test}) \left(p(x_l) = \begin{cases} 1, & \text{if } \langle d(x_l) \rangle_{predicted} = \langle d(x_l) \rangle_{observed} \\ 0, & otherwise \end{cases} \right)$$

where $\langle d(x_l) \rangle_{predicted, (s)} (s \in S_{training}) = v_d, v_d \in V_d$

DEFINITION 4. EFFICACY CRITERION OF PREDICTION

$$E1 = \frac{\sum_{l=1}^{|U_{test}|} p(x_l)}{|U_{test}|} \geq E1_{limit} > 0.5 \quad (\text{The efficacy of prediction is required to be better than random guess.})$$

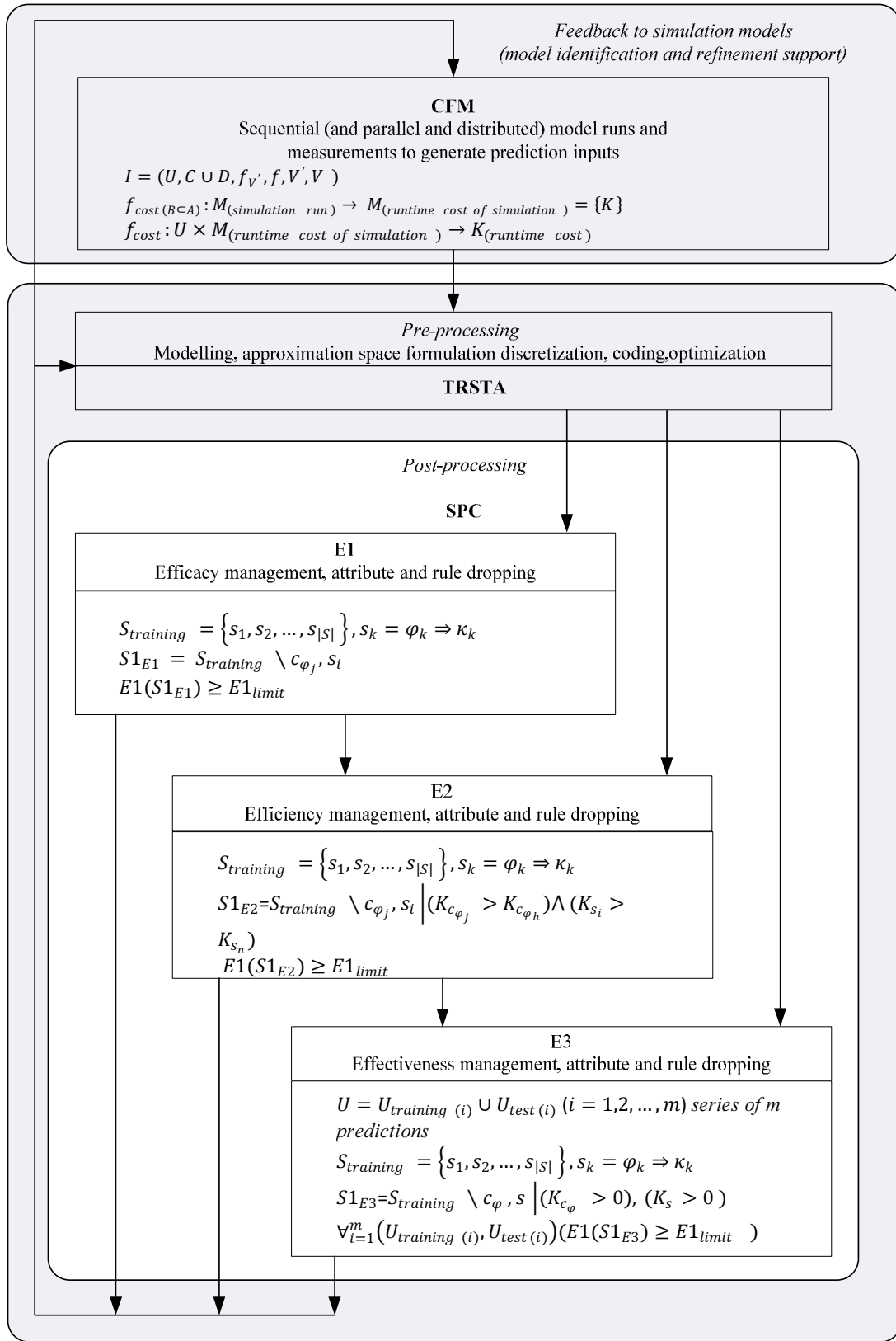


Figure 2. The Algorithm of Improved Simulation Performance Prediction (AISPP)

DEFINITION 5. THE COST ALLOCATION FUNCTION

The f_{cost} allocates the runtime costs of simulation in the information system

$$f_{cost}: U \times M_{(runtime\ cost\ of\ simulation)} \rightarrow K_{(runtime\ cost)}$$

DEFINITION 6. COST OF ATTRIBUTES, EXPERIMENTS AND RULES

Cost of an attribute a_i is defined as $K_{a_i} = \sum_{l=1}^{|U|} K_{a_i}(x_l)$ [sec], ($a_i \in A$), (cost of an attribute may also take the value of $K_{a_i} = 0$ [sec]).

In a decision table, the cost of an experiment x_l is determined as

$$K_{x_l} = \sum_{i=1}^{|C|} K_{x_l}(c_i)_{x_l} + K_{x_l}(d)_{x_l} \text{ [sec]}, (x_l \in U).$$

Cost of a rule s_i is calculated as

$$K_{s_i} = \sum_{\forall(x_l)|Match(s_i)=1} K_{x_l} \text{ [sec]}, (s_i \in S_{training}).$$

DEFINITION 6. EFFICACIOUS DROPPING OF ATTRIBUTES AND RULES

$$S_{training} = \{s_1, s_2, \dots, s_{|S|}\}, s_k = \varphi_k \Rightarrow \kappa_k$$

$$S1_{E1} = S_{training} \setminus c_{\varphi_j}, s_i$$

$$E1(S1_{E1}) \geq E1_{limit}$$

DEFINITION 7. EFFICIENT DROPPING OF ATTRIBUTES AND RULES

$$S_{training} = \{s_1, s_2, \dots, s_{|S|}\}, s_k = \varphi_k \Rightarrow \kappa_k$$

$$S1_{E2} = S_{training} \setminus c_{\varphi_j}, s_i \mid (K_{c_{\varphi_j}} > K_{c_{\varphi_h}}) \wedge (K_{s_i} > K_{s_n})$$

$$E1(S1_{E2}) \geq E1_{limit}$$

DEFINITION 8. EFFECTIVE DROPPING OF ATTRIBUTES AND RULES

$U = U_{training(i)} \cup U_{test(i)}$ ($i = 1, 2, \dots, m$) series of m predictions

$$S_{training} = \{s_1, s_2, \dots, s_{|S|}\}, s_k = \varphi_k \Rightarrow \kappa_k$$

$$S1_{E3} = S_{training} \setminus c_{\varphi}, s \mid (K_{c_{\varphi}} > 0), (K_s > 0)$$

$$\forall_{i=1}^m (U_{training(i)}, U_{test(i)}) (E1(S1_{E3}) \geq E1_{limit})$$

B. The AISPP process

Figure 2 shows the process diagram of the *new RST-based method – the Algorithm of Improved Simulation Performance Prediction (AISPP)*.

The *tactic* of the RST-based prediction improvement method can be formulated as follows: (1) all the data – got from sequential simulation and PADS (with different parameters – processor number, λ , etc.) – have to be analyzed which supposed to have influence on CFM prediction results. (2) CFM data, attributes and objects of an RST decision information system, are processed in a train-and-test examination process. (3) Using the defined SPC based evaluation of prediction performance criteria in the TRST train-and-test analysis, a feedback is realized to data collection and measurement and modelling steps of

simulation performance prediction (to model identification and refinement stage too).

AISPP process has the following features:

- the operations are applied with a traditional RST analysis method (TRST) and a traditional simulation performance prediction method (CFM)
- for the predictions, both the sequential and parallel and/or distributed simulation runs (if any) are used in the RST model
- the method functions in interactive manner using a TRST
- the *operations* that have been defined are for the analysis of the *efficacy* (E1), *efficiency* (E2) and *effectiveness* (E3) in the TRST process
- the method supports making *feedback* for model feature *identification and refinement support* of the simulation performance prediction models (or to the simulation model)

(The method can be implemented by using the OMNet++ DES software [12] and the ROSETTA Rough Set Software System [13].)

IV. CONCLUSIONS

For the improvement of the simulation performance prediction of a *parallel and/or distributed simulation model execution*, a new methodology, based on RST approach in order to work with unreliable, imprecise preliminary data, has been introduced.

For the new methodology, the set of necessary operations has been created:

- Operations, based on Systems Performance Criteria (SPC) of *efficacy* (E1), *efficiency* (E2) and *effectiveness* (E3) of predictions
- Operations for efficacious and efficient predictions attribute and rule dropping in predictions (for efficiency evaluation, the cost of attributes and cost of rules have been defined) and the effective attribute and rule droppings for a series of predictions too.

The algorithm of the new methodology has also been presented (Algorithm of Improved Simulation performance Prediction (AISPP)) with its connections to a traditional RST method (TRSTT) including approximation space optimization in pre-processing phase and embedding SPC in the post-processing. The new methodology provides feedback to the simulation performance model (and to the simulation model too) to support model features identification and refinement.

REFERENCES

- [1] Perumalla, K. S., 2010. “ $\mu\pi$: a scalable and transparent system for simulating MPI programs”, In Proceedings of the 3rd International ICST Conference on Simulation Tools and Techniques, ICST (Institute for Computer Sciences, Social-Informatics and Telecommunications Engineering), pp. 62.
- [2] Perumalla, K. S., 2006. “Parallel and distributed simulation: traditional techniques and recent advances”, In Proceedings of the 38th Winter Simulation Conference, pp. 84-95.

- [3] D'Angelo, G., 2011. "Parallel and Distributed Simulation from Many Cores to the Public Cloud (Extended Version)", In: Proceedings of the 2011 International Conference on High Performance Computing and Simulation (HPCS 2011), Istanbul (Turkey), IEEE, July 2011. ISBN 978-1-61284-382-7. <http://dx.doi.org/10.1109/HPCSim.2011.5999802> arXiv preprint arXiv:1105.2301. pp.1-9.
- [4] Yoginath, S. B.; Perumalla, K. S., 2013. "Empirical evaluation of conservative and optimistic discrete event execution on Cloud and VM platforms", In Proceedings of the 2013 ACM SIGSIM conference on Principles of advanced discrete simulation, pp. 201-210.
- [5] Fujimoto, R. M., 1990b. "Performance of Time Warp under Synthetic Workloads," in Proceedings of the SCS Multiconference on Distributed Simulation, San Diego, CA, USA, 17-19 January, 1990, pp. 23-28.
- [6] Fujimoto, R.M., 2000. "Parallel and Distributed Simulation Systems", John Wiley & Sons, USA, ISBN: 0-471-18383-0
- [7] Kunz, G.; Tenbusch, S.; Gross, J.; Wehrle, K. 2011. "Predicting Runtime Performance Bounds of Expanded Parallel Discrete Event Simulations", In Modeling, Analysis & Simulation of Computer and Telecommunication Systems (MASCOTS), 2011 IEEE 19th International Symposium on. IEEE, pp. 359-368.
- [8] Juhasz, Z.; Turner, S.; Kuntner, K.; Gerzson, M., 2003. "A Performance Analyser and Prediction Tool for Parallel Discrete Event Simulation", International Journal of Simulation, vol. 4, no. 1, May 2003, pp. 7-22.
- [9] Muka, L.; Lencse, G., 2008. "Cooperating Modelling Methods for Performance Evaluation of Interconnected Infocommunication and Business Process Systems", In Proceedings of the 2008 European Simulation and Modelling Conference (ESM2008), (Le Havre, France, Oct. 27-29.) EUROSIS-ETI, pp. 404-411.
- [10] Muka, L.; Lencse, G., 2011. "Method for Improving the Efficiency of Simulation of ICT and BP Systems by Using Fast and Detailed Models", Acta Technica Jaurinensis, Vol.5 No. 1. 2012., pp. 31-42.
- [11] Muka, L.; Benko, B. K., 2013. "Meta-level performance management of simulation: The problem context retrieval approach", Periodica Polytechnica, Electrical Engineering and Computer Science, 55(1-2), pp. 53-64, DOI: 10.3311/pp.ee.2011-1-2.06
- [12] Varga, A.; Hornig, R., 2008. "An overview of the OMNeT++ simulation environment", In Proceedings of the 1st international conference on Simulation tools and techniques for communications, networks and systems & workshops, ICST (Institute for Computer Sciences, Social-Informatics and Telecommunications Engineering), p. 60.
- [13] Komorowski, J.; Øhrn, A.; Skowron, A., 2002. "The ROSETTA Rough Set Software System", In Handbook of Data Mining and Knowledge Discovery, W. Klösgen and J. Zytkow (eds.), Oxford University Press, ch. 2-3, ISBN 0-19-511831-6.
- [14] Varga, A.; Sekercioglu, Y. A.; Egan, G. K. 2003. "A practical efficiency criterion for the null message algorithm", Proceedings of the European Simulation Symposium (ESS 2003), (Oct. 26-29, 2003, Delft, The Netherlands) SCS International, pp. 81-92.
- [15] Lencse, G.; Varga, A., 2010. "Performance Prediction of Conservative Parallel Discrete Event Simulation", Proceedings of the 2010 Industrial Simulation Conference (ISC'2010) (Budapest, Hungary, 7-9, June, 2010.) EUROSIS-ETI, pp. 214-219.
- [16] Lencse, G.; Derka, I.; Muka, L., 2013. "Towards the efficient simulation of telecommunication systems in heterogeneous execution environments", In proceedings of: TSP2013, Edited by Norbert Herencsar, Karol Molnar, the 36th International Conference on Telecommunications and Signal Processing, Rome, Italy, pp. 304-310.
- [17] Lencse, G.; Derka, I., 2013. "Testing the Speed-up of Parallel Discrete Event Simulation in Heterogeneous Execution Environments", In proceeding of: ISC'2013, Edited by Veronique Limere and El-Houssaine Aghezzaf, ISBN 978-90-77381-76-2, 11th Annual Industrial Simulation Conference, Ghent University, Ghent, Belgium, pp.101-107.
- [18] Lencse, G., 1998. "Efficient Parallel Simulation with the Statistical Synchronization Method", Proceedings of the Communication Networks and Distributed Systems Modeling and Simulation (CNDS'98) (San Diego, CA. Jan. 11-14). SCS International, pp. 3-8.
- [19] Lencse, G., 1999a. "Applicability Criteria of the Statistical Synchronization Method", Proceedings of the Communication Networks and Distributed Systems Modeling and Simulation (CNDS'99), San Francisco, CA. (1999. Jan. 17-20) SCS International, pp. 159-164.
- [20] Lencse, G., 1999b. "Design Criterion for the Statistics Exchange Control Algorithm used in the Statistical Synchronization Method", Proceedings of the Advanced Simulation Technologies Conference (ASTC 1999) part of the 32nd Annual Simulation Symposium (San Diego, CA. April 11-15) SCS International, pp. 138-144.
- [21] Pawlak, Z., 1998. "Rough set theory and its applications to data analysis", Cybernetics & Systems, 29(7), pp. 661-688.
- [22] Krishnaswamy, S.; Zaslavsky A.; Loke, S. W., 2002. "Predicting run times of applications using rough sets", In Ninth International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems (IPMU 2002), pp. 455-462.
- [23] Song, C.; Guan, X.; Zhao, Q.; Ho, Y. C., 2005. "Machine learning approach for determining feasible plans of a remanufacturing system", Automation Science and Engineering, IEEE Transactions on, 2(3), pp. 262-275.
- [24] Muka, L.; Derka, I., 2013. "Improving Performance Prediction of Parallel and Distributed Discrete Event Simulation: A Rough Sets-based Approach", In proceeding of: ISC'2013, Edited by Veronique Limere and El-Houssaine Aghezzaf, ISBN 978-90-77381-76-2, 11th Annual Industrial Simulation Conference, Ghent University, Ghent, Belgium, pp. 95-100.
- [25] Pawlak, Z.; 1982. "Rough sets", International Journal of Parallel Programming, 11(5), pp. 341-356.
- [26] Pawlak, Z.; Skowron, A., 2007. "Rudiments of rough sets", Information sciences, 177.1, pp. 3-27.
- [27] Zhang, J.; Li, T.; Ruan, D.; Gao, Z.; Zhao, C., 2012. "A parallel method for computing rough set approximations", Information Sciences, 194.2, pp. 209-223.
- [28] Bazan, J. G.; Szczuka, M., 2005. "The rough set exploration system", In Transactions on Rough Sets III, Springer Berlin Heidelberg, pp. 37-56.
- [29] Suraj, Z., 2004. "An Introduction to Rough Set Theory and Its Applications", ICENCO, Cairo, Egypt, pp. 1-39.
- [30] Gregory, F., 1993. "Cause, effect, efficiency and soft systems models", Journal of the Operational Research Society, vol. 44 (4), pp. 333-344.
- [31] Checkland, P., 2000. "Soft systems methodology: a thirty year retrospective", Systems Research and Behavioral Science, 17, S11-S58.
- [32] Muka, L.; Derka, I., 2014. "Simulation Performance Prediction in Clouds", In: Orosz Gábor Tamás (szerk.), 9th International Symposium on Applied Informatics and Related Areas – AIS2014., Konferencia helye, ideje: Székesfehérvár, Magyarország, 2014.11.12 Székesfehérvár, Óbudai Egyetem, pp. 142-147., ISBN:978-615-5460-21-0 **Error! Reference source not found.**

BIOGRAPHIES

LÁSZLÓ MUKA graduated in electrical engineering at the Technical University of Lvov in 1976. He got his special engineering degree in digital electronics at the Technical University of Budapest in 1981, and became a university doctor in architectures of CAD systems in 1987. Mr. Muka finished an MBA at Brunel University of London in 1996. Since 1996 he has been working in the area of modeling and simulation of infocommunication systems, which may include human subsystems too. Mr. Muka got his PhD at the Budapest University of Technology and Economics in 2011. Dr. Muka is working as an Associate Professor for the Széchenyi István University.

ISTVÁN DERKA received his Msc at Faculty of Electrical Engineering and Informatics at the Technical University of Budapest in 1995. He worked for the Department of Informatics from 1999 to 2003 and since then has been working for Department of Telecommunications, Széchenyi István University in Győr. He teaches Programming of communication systems and Interactive TV systems. He is an Assistant Professor. The area of his research includes multicast routing protocols and IPTV services in large scale networks.

Determining the Aging of Cheese from Sheep's Milk by "Electronic Nose" Multi-sensor System

Todor Todorov*, Stefan Ivanov*, Toshko Nenov* and Marta Seebauer**

* Technical University of Gabrovo, Gabrovo, Bulgaria

** Óbuda University/Alba Regia Technical Faculty, Székesfehérvár, Hungary
t61@abv.bg, st_ivanov@abv.bg, tgnenov@gmail.com, seebauer.marta@amk.uni-obuda.hu

Abstract— Production of a high quality cheese is a complex and time consuming process requiring exact compliance with the technology of manufacturing. This work presents the results of test of sheep's milk cheese done by "electronic nose" multi-sensor system based on metal oxide gas sensors. Several samples of cheese were measured during different stages of its aging. The collected data is used for training, validation and test of artificial neural network for determining the aging of sheep cheese.

I. INTRODUCTION

Rapid quality control of foods and the determination of their quality during storage time are essential for human health.

For determining of their taste usually is used organoleptic sensory analysis called "flavor analysis." For this purpose is employed a team of well-trained tasters. The essence of tasting process in practice is that the tasters assess the product through their receptors for taste, odor, smell and vision [1]. The disadvantages of this method of assessment are that it is relatively expensive, requires a lot of time and has a rather subjective character.

Another method to diagnose the state of foods is the determining the composition of gases collected from the product by using a classical gas chromatograph [2]. This method is considerably objective and accurate but is characterized by its high costs and long duration of measurements.

The most accurate of all analysis is the method based on examination of the specific protein composition and DNA identification [3]. In this case, like in the first two methods also, serious disadvantages are the high costs, and the need for relatively long time for testing, as well as the requirement for highly trained personnel, and expensive special equipment.

Knowing the working principles of the olfactory organs of biological systems allows the development of technical tools for gas analysis based on gas sensors which are called "electronic nose" and are comparable in efficiency with their biological analogues [4]. The "electronic nose" is multi-sensor system for detection of gaseous substances. Different from traditional sensor systems requiring highly selective sensitive elements, the "electronic nose" uses a set of relatively non-selective gas sensors. The development of "electronic nose" systems is based on the modern technologies in the fields of electronics and special methods for data processing and data fusion [5].

The "electronic nose" is used in the food industry for quality control, monitoring and evaluation of the products freshness, determining of their expiration date and others.

Its advantage is based on the fact that it provides the ability to determine the quality of the tested products quickly and with great accuracy [6].

In this work are presented the testing results of cheese from sheep's milk carried out with "electronic nose" multi-sensor system based on metal oxide gas sensors. Tested samples of cheese were taken during different stages of its aging. Artificial neural network was trained and tested with the collected data to determine the maturity of the sheep cheese.

II. EXPERIMENTAL RESEARCHES AND RESULTS

The metal-oxide gas sensors are particularly suitable for the realization of the "electronic nose" because they have a high sensitivity related to different gases and - in addition to this - they have a very low price. They are used in systems for detection of intentionally added water in milk [7], for analysis and quality control of yoghurt [8], for detection of dangerous bacteria in dairy products [9, 10], etc.

In the experiments was used cheese from sheep's milk. Samples of sheep cheese were obtained directly in production plant and they were taken during the overall period of its aging. The different samples for experiments were taken from 5th, 14th, 21st, 28th and 46th day of the cheese production. The analysis of individual samples was carried out by multi-sensor system of type "electronic nose" based on metal oxide gas sensors as described in [11].

All tested samples for a given period of aging were divided into equal parts in shape and weight and stored in a refrigerator with temperature at about 4°C. Before the measurement of each sample, the "electronic nose" system was turned on for a specified time, to reach by the test chamber and the gas sensor the operating temperature. During the process of measurement the glass container, in which the sample is placed, is maintained at constant temperature of 50°C. Placing the sample in an environment with such a temperature contributes to the separation of a larger quantity of gases which are analyzed by the gas sensors in the upper part of the chamber. The period of measurement for every individual sample is equal. After gathering the results of all sensors and removal of sample it is required a definite time for relaxation of sensors. In the study for every period of aging were carried out six measurements with different samples. Five of the results were used for training of artificial neural network. The results of the rest measurement were used for testing of the trained network. Samples used for every testing were selected by random way.

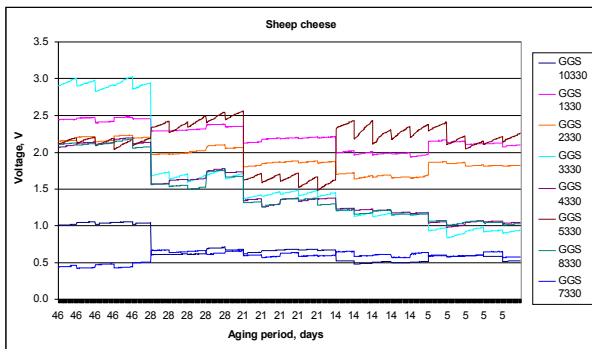


Figure 1. Signal levels from sensors for five measurements of sheep cheese

The data that were used to train the neural network are shown in the following figures.

Fig.1 shows the data from five measurements on five different aging periods of sheep cheese.

In the results presented in Fig.1 can be observed the differences in the reactions of six sensors (GGS 1330, GGS 2330, GGS 3330, GGS 4330, GGS5330 and GGS 7330) according to different periods of aging of sheep cheese. The results show that there is a good repeatability in the responses of sensors for the five subsequent measurements of cheese samples.

Fig.2 presents the data about the averaged measurements of each cheese sample. The figure clearly shows the trend of changing responses of the sensors. The most informative in this case are the sensors GGS 3330 and GGS 4330. These two sensors demonstrate the best expressed changes in the response of sensors to the period of aging of the cheese.

Fig.3 shows the averaged data for each of the five different periods of aging of sheep cheese.

When the data of five measurements for each stage of aging are averaged, the responses from the gas sensors have very definite changes for the different periods. From the results it can be concluded that when signals from five sensors (GGS 1330, GGS 2330, GGS 3330, GGS4330, and GGS 8330) are observed then for good separation for different periods of aging of the cheese can be marked. The best separation for the different periods of aging was produced by the signals of three sensors (GGS 3330, GGS4330 and GGS 8330).

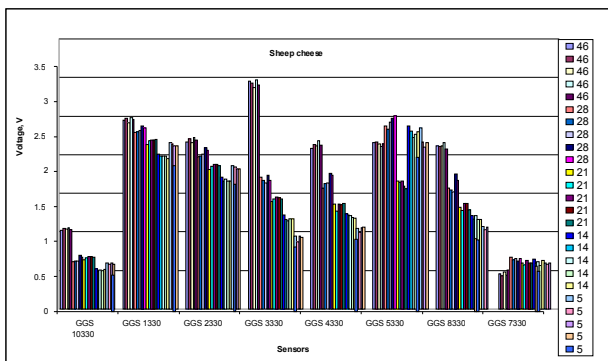


Figure 2. Averaged values of signals from sensors for measurement of five samples of sheep cheese

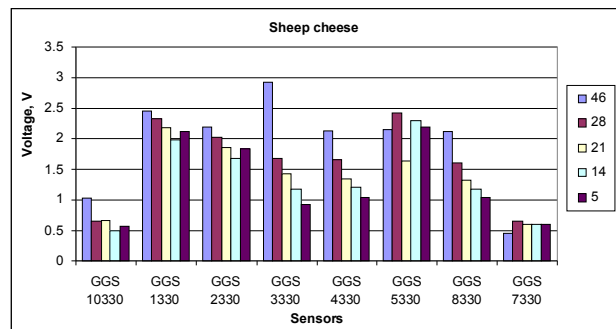


Figure 3. Averaged values from the signals of sensors for the five measurements of each period of aging of sheep cheese

The trend in the responses of the sensors is very clear. The differences in the levels of the sensors are directly related to the period of aging of the cheese.

III. ARTIFICIAL NEURAL NETWORK FOR ANALYZE OF AGING OF CHEESE FROM SHEEP'S MILK

To determine the quality of sheep cheese in accordance with the level of its aging was trained an artificial neural network. The Levenberg-Marquardt algorithm was used as training algorithm. The structure of the neural network consists of 3 inputs, one hidden layer with 10 neurons and one neuron in the output layer. The transfer function of the neurons in hidden layer is tangent sigmoid function and a linear function is used for the neuron in the output layer.

The data for training of neural network were five of the six samples of sheep's cheese with varying degrees of aging. After the data obtained from all sensors were analyzed, only data from three sensors - GGS3330, GGS4330 and GGS8330 were selected for training of the neural network while these three sensors had the best reaction depending on the degree of aging of sheep cheese.

Before the training of a neural network, a polynomial approximation was used to generate additional points that lie on approximated data curves. A second order polynomial process was used for approximation. Fig.4 shows the dependences for the three sensors which were received by the approximation.

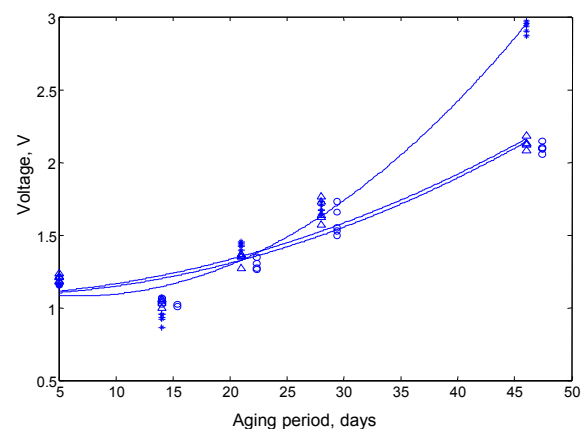


Figure 4. Approximated data from three sensors for determining of the time of aging of sheep cheese

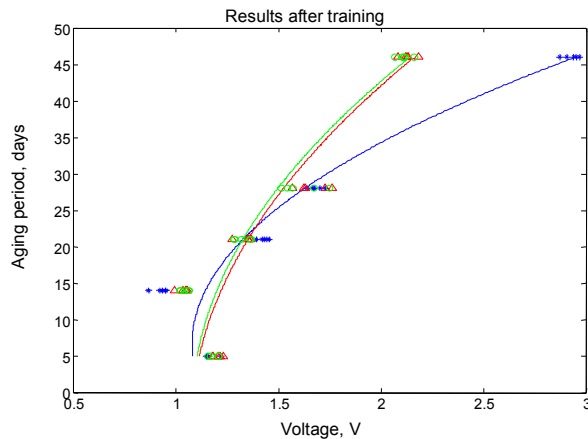


Figure 5. Results of a neural network training for determining of the time of aging of sheep cheese

After the training, the performance of the neural network has been verified with testing data. The test results are presented on Fig.5.

The comparison between the expected time and the time received on the neural network output (Fig.6) it can be seen that the neural network was trained successfully and there is no big discrepancy between the expected and received time. The error that occurs in this case is about 10 minutes, which is negligible for a period of aging which is measured in days and weeks.

The work of the trained neural network was tested with input data which has random noise added with maximum amplitude of 50mV.

Fig.7 presents the relation between the expected time of aging and the resulting time generated by the neural network with presence of noise in the input data.

The proper functioning of neural network except with the simulated data was checked with testing data, which have not participated in the training process of the network.

When the neural network was tested its response determined successfully submitted data for different aging periods. Fig.8 illustrates the distribution of neural networks response to the testing data in the three-dimensional space using for visualization data from two sensors.

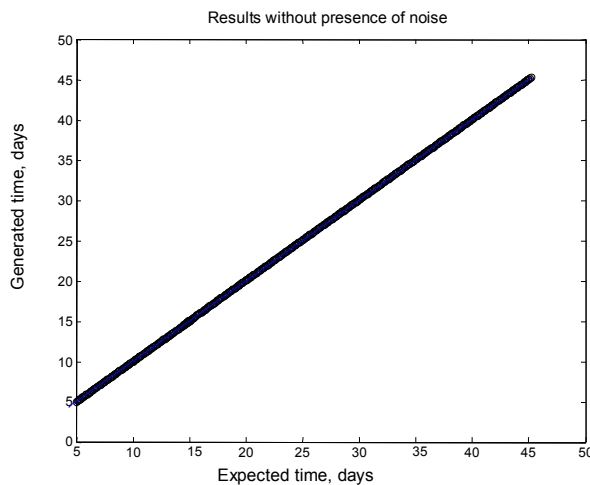


Figure 6. Time - expected and received from the neural network for the period of aging of sheep cheese

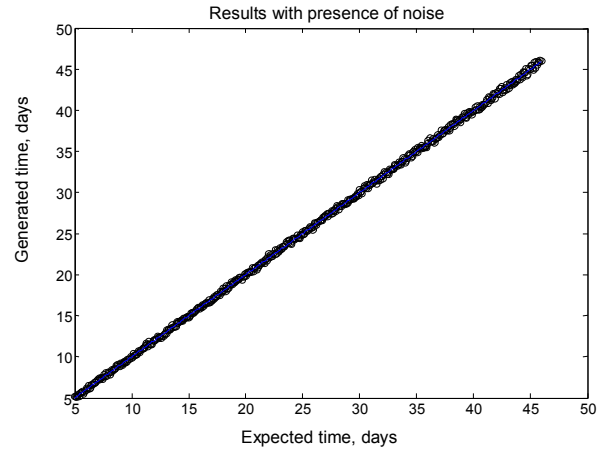


Figure 7. Expected and received from the neural network time for the period of aging of sheep cheese with the presence of noise in the input data

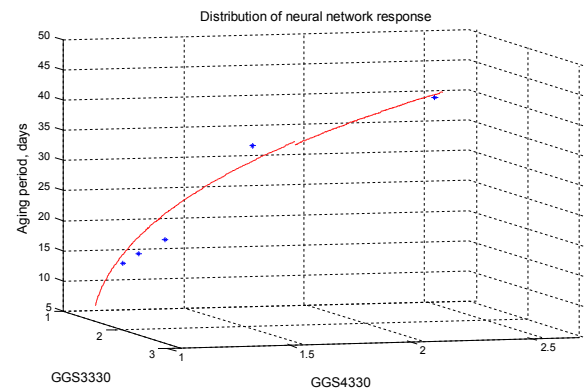


Figure 8. Distribution of testing data for the time of aging of sheep cheese

Fig.9 illustrates the time of aging of sheep cheese generated by the neural network with actual testing data. The straight line presents the expected time in days, the other line – the time derived from neural network with applied testing data. The biggest is the deviation in the initial period of aging - 6 days. Increasing the time of aging of the cheese, the deviation decreases and at 46 day the deviation in aging is within one day, which is sufficient to determine the state of the cheese.

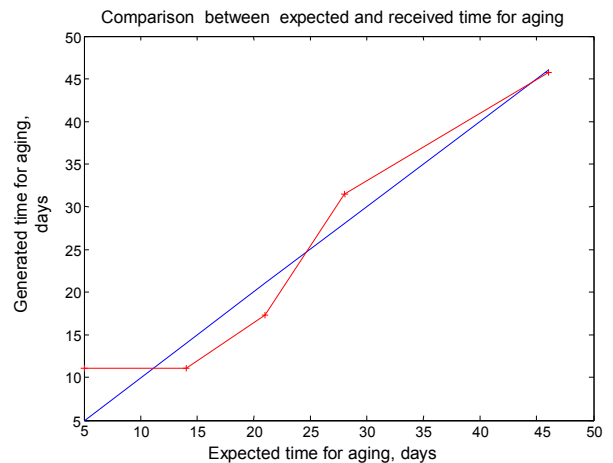


Figure 9. Time received by the neural network for the period of aging of sheep cheese

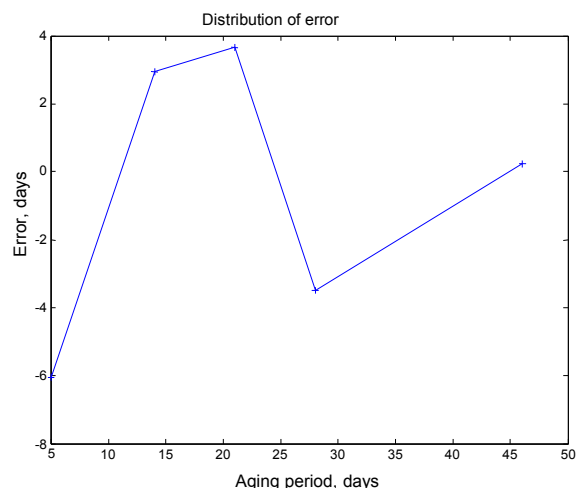


Figure 10. Distribution of error for generated by the neural network results for the period of aging

Fig.10 presents the error for time of aging of sheep cheese achieved from neural network with the actual testing data.

IV. CONCLUSION

The work presents the experimental test of cheese from sheep's milk. The analyzed samples of cheese were taken during different stages of its aging. The analyses were done with the help of multi-sensor system for quality control of food products based on metal oxide gas sensors. Results of this study show significant differences in the responses of gas sensors. They are in direct relation to the period of aging of the various samples of the cheese.

With the data obtained from the tests was trained an artificial neural network. When the neural network was tested with data that have not participated in the training it was found that the error decreases with time and at full aging it is within the range of one day. The cheese has to be qualified at the final stage of its aging and therefore an error of one day is normally and allowable when recognition by neural network is implemented.

Based on the obtained results it can be concluded, that a multi-sensor system could be applied successfully for analyze and determination the aging of cheese from sheep's milk which is commercially available.

REFERENCES

- [1] G. Zsivanovits, D. Iserliyska, Haracterization of the texture of food products, *New Knowledge Journal of Science*, 2013, pp. 111-115.
- [2] B. Kolb, L. Ettre, Static headspace-gas chromatography, John Wiley&Sons, Inc., Hoboken, New Jersey, 2006.
- [3] D. Kalaidjiev, Genetic and average variability of the coagulation ability of milk, PhD Thesis, Agricultural Academy, Agricultural Institute, Stara Zagora, 2014 (in Bulgarian).
- [4] P. Barlett, J.Elliott, J.Gardner, Electronic Noise and their application in the food industry, *Food Technology*, 1997, vol. 51, No. 12, pp. 248-256.
- [5] T. Pearce, S.Schiffman, H.Nagle, J.Gardner, *Handbook of machine olfaction*, Wiley/vch Verlag GmbH & co. Weinheim, UK, 2003, pp. 138.
- [6] H. Swatland, Effect of connective tissue on the shape of reflectance spectra obtained with a fibre-optic fat-depth probe in beef. *Meat Sci.* 2001, vol. 57, pp.209-213.
- [7] H. Yu, J.Wang, Y.Xu. Identification of Adulterated Milk Using Electronic Nose. *Sens. Mater.*, 2007, vol. 19, pp. 275-285.
- [8] G. Green, A.Chan, R.Goubran. Tracking Food Spoilage in the Smart Home Using Odour Monitoring. in *Proceedings of the 2011 IEEE International Workshop on Medical Measurements and Applications (MeMeA)*, Bari, Italy, 30-31 May 2011, pp. 284-287.
- [9] Z. Ali, W. T. O'Hare, B. Theaker. Detection Of Bacterial Contaminated Milk By Means Of A Quartz Crystal Microbalance Based Electronic Nose. *Journal of Thermal Analysis and Calorimetry*, 2003, vol. 71, pp.155-161.
- [10] S. Benedetti, F.Bonomi, S.Iametti, S. Mannino, M.Cosio. Detection of aflatoxin M1 in ewe milk by using an EN. In *Proceedings of the 2nd Central European Meeting 5th Croatian Congress of FTBN*, Opatija, Croatia, October 17-20, 2004, pp. 101-105.
- [11] T. Todorov, T. Nenov, S. Ivanov. Multisensor system for quality control of food products, *Automatics and Informatics*. 2014, No 3, pp.41-46 (in Bulgarian).

IT Promoting Physics Projects

C. Fülöp*, R. Szabó, T. Berényi and B. Simó

Trefort Ágoston Bilingual Technical High School, Budapest, Hungary

*ELTE TTK Physics Education PhD Program, Budapest, Hungary

fulcsilla@gmail.com, hu.szabo.roland@gmail.com, berenyi.tamas444@gmail.com, slmob4lu@gmail.com

Abstract- STEM (S-Science, T-Technology, E-Engineering, M-Mathematics) is in crisis. Research in didactics is done to find a way out. In physics education research IT plays an important role. Students are fond of new technology, and are familiar with informatics. Surely, cutting edge science research is in need of using IT. Obviously, teaching new contents in most topics can't miss out using some kind of, but simplified tools or programs.

IT can call to life projects in the field of classical physics, too. In our project we analyzed a simple motion on a slope the classical way, which led us to a function of two variables. For a numerical analysis we needed computing skills. Our group achieved surprising results in the study.

Key words: classical physics, mentor class, numerical analysis in high school

I. INTRODUCTION

We have to be aware of the fact that STEM (S-Science, T-Technology, E-Engineering, M-Mathematics) is in crisis. Students don't like these areas, and less and less choose a profession from it as carrier.

Research is done in didactics how to make these fields more enchanting and understandable for the future generation.

Using applied or high-tech apparatus, or combining information technology or computing in our projects helps to attract students to a decision to participate according to the experience.

In physics education research one of the biggest tasks is to develop projects and methodological tools for new content. One of the goals is to teach the latest achievements of science research. As IT is an essential tool in the research itself, surely it plays a determinative role in the education, too. Particle physics, environmental science chaos theory and similar content have to build on IT even in the popularization of these topics. If we are talking about special classes or projects for those who have special talent or pay interest in science (we call these mentor classes or mentored projects), even more of IT we need to rely on.

We present "The sledge project" (Ref. [1].) as an example of a partly different project from what we discussed above. We made a deeper study of a question from the fields of classical physics. After a few steps we got to a function of two variables. Analyzing it is not in the secondary school curriculum. But programming made it possible to answer our questions.

A. Topic flag

Travelling home from a physics competition with eager students, we discussed some issues that arose during the event. A well-known question came into our focus:

"Why is it easier to pull a sledge horizontally than to pull it on a slope upwards?" (See Fig. 1.)



FIGURE 1.
PULLING A SLEDGE

Some of the students paid special attention to the problem and showed interest to participate in a deeper study.

B. A study of the typical qualitative answers

We met answers that are similar to this reasoning: "In both cases we need to exert a force against friction, plus on the slope we must exert force against a component of gravity as well."

The problem is as follows: We agree that the force against gravity is an increasing value as the tilt angle of the slope is increasing, whereas the force against friction is a decreasing one. When we add up these two, the sum is not necessarily an increasing function.

Some of the answers we know of are similar to this: "Besides energy dissipated in friction, extra mechanical work must be done to give „height”/ „positional”/ „potential”/ „gravitational” energy."

When talking about an "easier pull" we associate it with forces rather than with energy. Work, energy, and force are different notions. If we want to make a connection, we need to study distance as well.

C. The Newtonian analysis

We used Newton's laws, which are also well known as basics of classical dynamics for the theoretical analysis of the case. We denote the notions used in the analysis in dynamics by the symbols used in SI system: F , m , a , μ , α , etc..

Based on Newton's 2nd law the force needed for a uniform motion...

- ... in case of pull on level ground (denoted by *) is as follows:

$$* F_{\text{pull}} = - * F_{\text{friction}} \quad (1)$$

$$\text{, since } \sum \mathbf{F} = 0, \quad (2)$$

$$\text{so } \mathbf{F}_{\text{pull}} = \mu \cdot m \cdot g \quad (3)$$

- ... in case of pulling up on a slope (see Fig. 2.) is as discussed below:

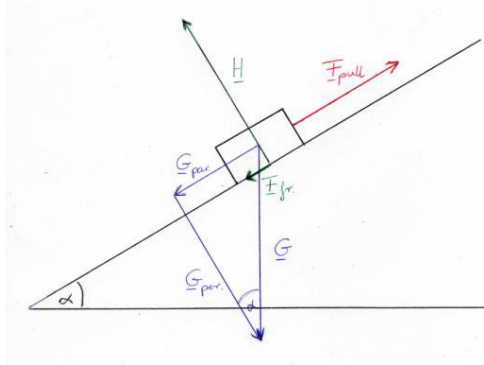


FIGURE 2.
A STUDY OF THE FORCES ON A SLOPE

- The force that is exerted by the slope is H.

$$\mathbf{H} = -\mathbf{G}_{\text{perpendicular}} \quad (4)$$

, therefore

$$H = m \cdot g \cdot \cos \alpha \quad (5)$$

- We can calculate friction, since

$$F_{\text{friction}} = \mu \cdot H \quad (6)$$

$$\text{, so } F_{\text{friction}} = \mu \cdot m \cdot g \cdot \cos \alpha \quad (7)$$

- The component of gravity can be given as

$$\mathbf{L} = -\mathbf{G}_{\text{parallel}} \quad (8)$$

, that means

$$L = m \cdot g \cdot \sin \alpha \quad (9)$$

Based on Newton's 2nd law the force of pull is

$$\mathbf{F}_{\text{pull}} - \mathbf{F}_{\text{friction}} - \mathbf{L} = \mathbf{0} \quad (10)$$

, which gives us that

$$F_{\text{pull}} = \mu \cdot m \cdot g \cdot \cos \alpha + m \cdot g \cdot \sin \alpha \quad (11)$$

$$F_{\text{pull}} = m \cdot g \cdot (\mu \cdot \cos \alpha + \sin \alpha) \quad (12)$$

To compare the force of pull in these cases we formed a function

$$\psi = F_{\text{pull}} - F_{\text{pull}} \quad (13)$$

We received that

$$\psi = m \cdot g \cdot (\mu \cdot \cos \alpha + \sin \alpha - \mu) \quad (14)$$

If we study the $\text{sgn}\psi$ function, we can figure out if our original statement is true or false. We need to face a problem: analysing a function like $\text{sgn}\psi$ is not in the secondary school curriculum. A few decades ago we could not have coped with a task like this analysis.

D. Numerical analysis, a study of the $\text{sgn}\psi$ function

Two of the students participating in the project are students of IT software.

So we wrote a programme in C++ using SDL (1000x180 pixels). Since $0^\circ \leq \alpha \leq 90^\circ$ on the vertical axis we can easily represent the tilt angle, α if $1^\circ = 2$ pixels. So, on the horizontal axis we can represent μ . With a multiplier we can adjust the maximum value to what we want to study. Our programme works in two cycles. This means 90,000 data-pairs to calculate with. We presented the results according to our purpose in a colour code (Table 1.)

TABLE 1.
COLOR CODE OF $\text{SGN}\Psi$ FUNCTION IN SDL

Pull on slope	Pull on level ground	$\text{sgn}\psi$	Colour code
bigger	smaller	+	red
smaller	bigger	-	blue

We were very excited to see the results. If blue area appears, it means that the original statement is not necessarily true in all circumstances. Our results in the numerical analysis:

- For all possible angles, if $0 \leq \mu \leq 0.25$, see Fig. 3.

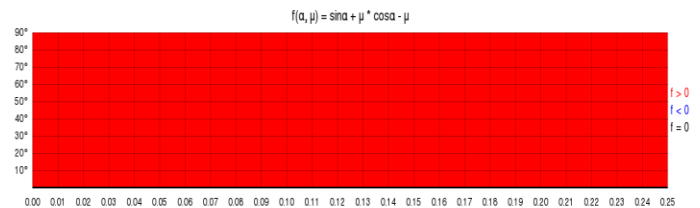


FIGURE 3.
 $\text{Sgn}\psi$, if $0 \leq \mu \leq 0.25$

- For all possible angles if $0 \leq \mu \leq 2.5$ we present the result in Fig. 4.

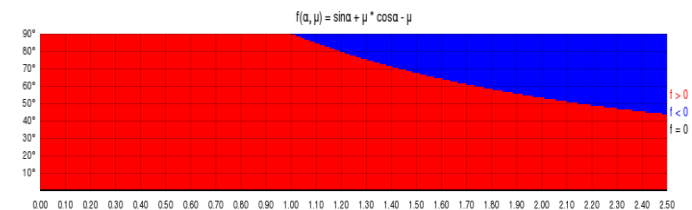


FIGURE 4.
 $\text{Sgn}\psi$, if $0 \leq \mu \leq 2.5$

The fact that a blue area appears on the figure means that the statement is not even necessarily true.

3.) For all possible angles, we allowed μ up to 50, you can check Fig. 5.

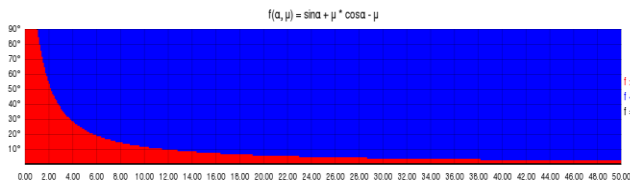


FIGURE 5.
Sgn ψ , if $0 \leq \mu \leq 50$

We can conclude that α and μ are the main quantities that define motion on a slope.

E. Hands-on" measurements

We wanted to see what the typical values (for μ and α) are, when playing the sledge.

• Measuring the friction constant

We pulled the sledge on level ground at constant speed. We used an 80213-141 Kamasaki digital scale (dynamometer) that we bought cheap in a fishing shop. We also needed a bathroom scale and a sledge. We made measurements 3 different occasions, which means 3 different circumstances. We decided to note 3 readings each time. We formed the mean value by calculating the arithmetic mean. Our results are in Table 2.

TABLE 2.
OUR RESULTS FOR FRICTION CONSTANT

measurement (Budapest XXI. ÁMK)	$F_{gravity} (N)$	$F_{pull} (N)$	$F = \frac{F_{pull}}{F_{gravity}}$	μ_{mean}
9 th Febr. 2015. late evening	351+51.7= 403	45.15	0.112	0.118
		49.46	0.123	
		47.88	0.119	
10 th Febr. 2015. afternoon	51.7	9.88	0.191	0.178
		9.20	0.178	
		9.45	0.166	
16 th Febr. 2015. early morning	51.7	4.90	0.095	0.092
		5.10	0.098	
		4.35	0.084	

In journal "KÖMAL" (, which is a mathematics and physics journal in Hungary for higher primary and secondary school students) we found that the friction constant measured with a different method is $0.02 \leq \mu \leq 0.3$. Our results match those we found in the literature (Ref.[2].).

➤ Measuring tilt angles

Our first problem was that we could not get an inclinometer. Since this instrument is not cheap, we worked out a conventional method for α . We needed a bubble level (0.8m), a 1meter rod, and a pendulum (string & load). Fig.6 and 7. demonstrate how we used our tools.

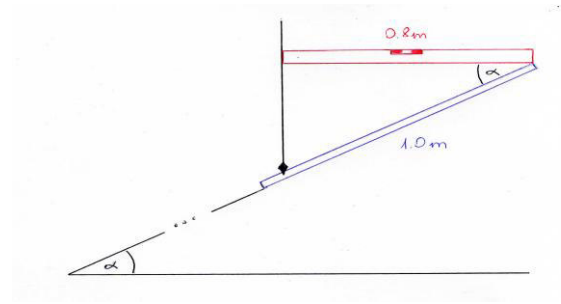


FIGURE 6.
OUR HAND-MADE APPARATUS

We also used applied apparatus to measure tilt angles, the GPS system. We worked with the two versions that were available free. See Fig. 7./b.



FIGURE 7.
IN-SITU MEASUREMENTS

We made our measurements in different playgrounds in Budapest (Bp.) on 23rd June 2015. We visited 3 playgrounds:

- Slope 1. is Bp. 1095 Petőfi utca 2.
- Slope 2. is Bp. 1091 Kékvirág utca 2.
- Slope 3. is Bp. 1107 Bihari utca 3-5.

Table 3. includes our results. We denote by * our results with the GPS system.

TABLE 3.
TILT ANGLES OF SLEDGING SLOPES

slope	spot	l_{pre} (cm)	$\cos \alpha$	α_{act}	α mean	* α act1	* α act2	* α mean
S1	1/1	84.0	0.9524	18°	15°	16°	13°	15°
	1/2	85.0	0.9512	20°				
	1/3	80.5	0.9938	6°				
S2	2/1	80.5	0.9938	6°	11°	11°	14°	12°
	2/2	81.5	0.9816	11°				
	2/3	83.3	0.9639	15°				
S3	3/1	83.5	0.9581	17°	17°	15°	14°	15°
	3/2	85.0	0.9412	20°				
	3/3	82.5	0.9697	14°				

Our results range from 6° to 20° , and the mean value is 14°.

F. Incorporating the results of our theoretical and practical studies

„Why is it easier to pull a sledge on level ground than to pull it up a slope?”

We provide two answers:

1.) Since $\mu < 1$, from the theoretical study we learned that there is no need to give a typical value to α . (Fig. 8.)

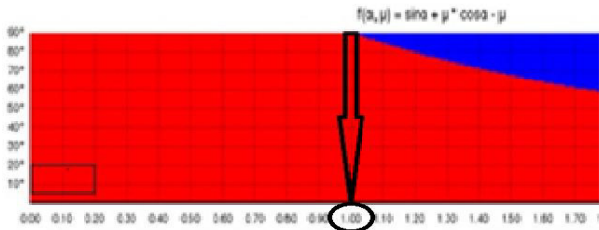


FIGURE 8.
WHEN THE FRICTION CONSTANT IS BELOW 1

A correct answer is: As the typical $\mu < 1$, it is easier to pull a sledge on level ground than to pull it up a slope.

2.) We studied the area denoted by the typical values based on our measurements (Fig. 9.).

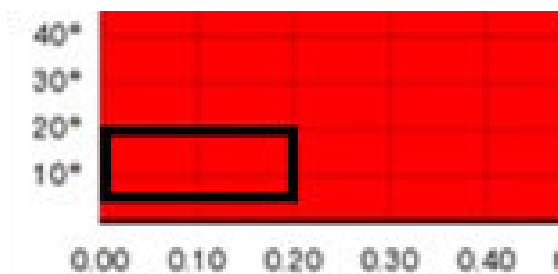


FIGURE 9.
REAL SLEDGING

Another correct answer is: It is easier to pull a sledge on level ground than to pull it up a slope, because of the real values of α and μ .

CONCLUSION

In science methodology revolution can break out as in research IT is becoming a new tool in education.

In “The Sledge project” we studied an often-asked question in classical mechanics. This topic has been studied before only in a qualitative way in primary and secondary school science education. We have found that the typical answers in this case are not correct. We faced a big problem: reasons were asked for, whereas the statement is not obvious, furthermore: the statement is not necessarily true in all physical conditions.

Easily we got to a function of two variables. From this point on computing skills were needed to carry on with the project. Our peripatetic mentor class turned into an interdisciplinary forum of science. We made quantitative analysis. We could find real answers for a classical physics question, where the qualitative answer based on experience opposed the scientific truth.

ACKNOWLEDGMENT

We thank the academic staff of Trefort Ágoston Bilingual Technical High School, Budapest for providing the physical conditions we needed. Special thanks to Károly Varsóci, the Headmaster of the school, and Péter Klettner, the teacher of IT for their support. Also, Csaba Dér, the school’s system administrator, and Otília Bartha, the Head of Science Department followed attentively and helped our research.

C. Fülöp expresses her gratitude to Judit Illy, her PhD consultant.

REFERENCES

- [1] C. Fülöp, R. Szabó, T. Berényi, B. Simó, “The Sledge Project”, unpublished
- [2] K.Tar, “77. Mérés Feladat” in KÖMAL, MATFUND, pp. 238-239., May 1985.

Relative Orientation in Photogrammetry

Tamas Jancso

*Institute of Geoinformatics, Alba Regia Technical Faculty, Óbuda University
Pirosalma u. 1-3, Székesfehérvár, H-8000, Hungary*

jancso.tamas@amk.uni-obuda.hu

Abstract— In photogrammetry to form a 3D model we need at least a stereo pair of images taken from two different stations. By this way we get two photos having their own perspective centres. In order to observe the stereo model we have to calculate the position of images in space in the moment of exposure. There are several solutions for this problem, although we use most frequently the relative orientation approach. By this task we try to calculate the position of images comparing to each other in a common spatial coordinate system. In other words we can form the task as a geometric problem: we are looking for the position of two planes (images) in the space and we have to describe the relation of these two planes with the minimally necessary parameters. Present paper summarizes the most common computational method emphasizing the problems and the limitations for applying it for photogrammetric tasks.

Keywords— Relative orientation, Stereo pair, Least square method

I. INTRODUCTION

A stereo pair is oriented relatively if the projection rays of corresponding points intersect each other in the model space. In other words the pair of rays should be on the same plane otherwise we experience a vertical parallax when observing the stereo model for example with anaglyph glasses. This condition is called coplanarity (Figure 1). If we want to describe the coplanarity condition we need a spatial coordinate system bound to our stereo model. In this coordinate system we have to describe the position of image planes. Present paper summarizes the different approaches for relative orientation. The common point in all methods that we need five parameters to determine the orientation of images.

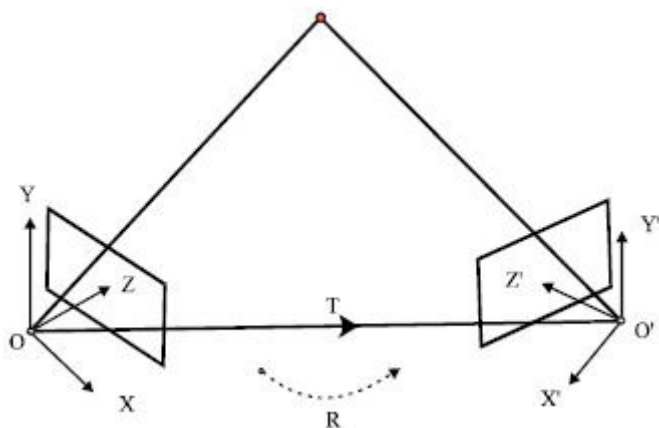


Fig. 1 Coplanarity condition

II. COORDINATE SYSTEMS AND PARAMETERS

The relative orientation actually is not a transformation problem, which means our goal is not to assure a pathway between different coordinate systems. Instead, we want to define the relative position of the two images in the model coordinate system to serve as a spatial reference system. All subsequent processing steps are performed in the model coordinate system, as the projection rays of common points are defined in this system as illustrated in Figure 2 [1].

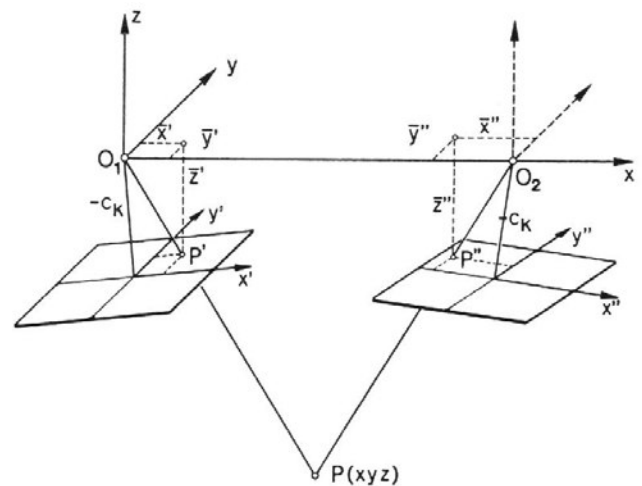


Fig. 2 Coordinate system for relative orientation

A. Different approaches

The question may arise whether it is necessary to rely on this coordinate system, to be more precise, is it necessary to know the model coordinates of a field point? In addition, we need to form the spatial coordinates $\bar{x}', \bar{y}', \bar{z}', \bar{x}'', \bar{y}'', \bar{z}''$ of image points P', P'' located in the image plane but considered as model points. It would seem much preferable, if the image coordinates x', y', x'', z'' could be used immediately to get the terrain coordinates. Indeed collinear equations create a direct link between the image and the terrain point.

The reason why the relative orientation is needed is the stereoscopic viewing and pointing, as a measurement method. Earlier in the era of analogue instruments the stereo-image pairs were considered as real planes in the space and they

were moved and rotated around three coordinate axis inside the evaluation instrument (Figure 3).

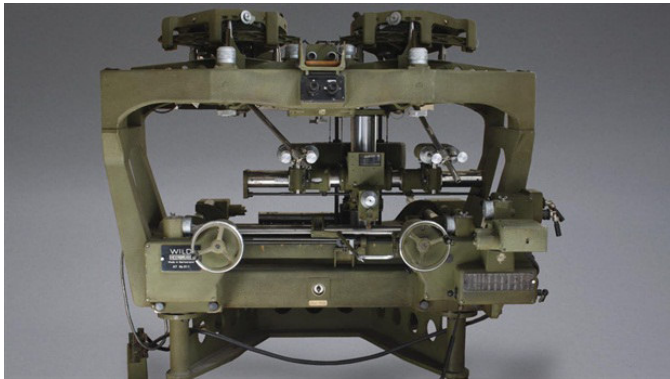


Fig. 3 Analogue photogrammetric evaluation instrument

The relative orientation is actually an attempt to restore the situation in which the images were situated relatively to one another at the moment of image recording. If we skip this step, the stereoscopic image observation would not be possible because of the conjunctive projection rays would not be in the same plane, which means they would not intersect each other.

Later there was a need for the stereoscopic observation even for the analytical stereo comparators and digital photogrammetric workstation; and therefore, the relative orientation is still part of the orientation process (Figure 4).



Fig. 4 Digital Photogrammetric Workstation (DPW)

However, the collinear equations mathematically and accurately describe the strict connection between the corresponding image and field points, so we prefer to calculate the field coordinates of points with this method. As a consequence, the model coordinates are no longer needed, as the terrain coordinates are directly calculated from image coordinates. In practical terms, the relative orientation should be carried out to perform the stereoscopic viewing. Of course, there are several digital evaluation programs which offer the conventional orientation routes, after which the field coordinates are calculated with the spatial similarity

transformation based on model coordinates, but it is actually not necessary and it is getting increasingly neglected.

Basically we distinguish two types of relative orientation, depending on the set of parameters. In the first variant just the right image is moved and rotated, in the second variant both images are rotated. In the digital photogrammetric evaluation both modes may be selected since the images are not rotated and shifted relative to each other in reality, the orientation elements are calculated only. The resulting orientation

elements can be used to calculate the theoretical y''^0 at a given p_x parallax. Next, the right image is moved in y -direction on the screen, while the measured image coordinate y'' within a certain margin of error does not correspond to the theoretical value of y''^0 , so we enable a smooth stereoscopic observation. The theoretical vertical parallax is calculated by:

$$y''^0 = \frac{\begin{vmatrix} b_x & b_y & b_z \\ \bar{x}''^0 & \bar{y}''^0 & \bar{z}''^0 \\ \bar{x}' & \bar{y}' & \bar{z}' \end{vmatrix}}{\begin{vmatrix} b_x & b_y & b_z \\ \bar{x}' & \bar{y}' & \bar{z}' \\ r''_{12} & r''_{22} & r''_{32} \end{vmatrix}}$$

Symbols:

b_x, b_y, b_z : Base vector components, in other words, the vector between $O_1 O_2$

$\bar{x}', \bar{y}', \bar{z}', \bar{x}'', \bar{y}'', \bar{z}''$: Model coordinates of image points P', P''

$r''_{12}, r''_{22}, r''_{32}$: Elements of the rotation matrix (direction cosines)

$\bar{x}''^0, \bar{y}''^0, \bar{z}''^0$: Theoretical model coordinates of P'' (we assume $y'' = 0$)

B. Solution and optimization

The initial conditions equation of the relative orientation is deduced from a simple but ingenious geometrical consideration. Intersecting rays of a model point are connected with four notable points in the model coordinate system. These points are O_1, O_2 as the perspective centres,

and the image points of P', P'' corresponding to the model point P . All these points are determined in the model coordinate system. (Figure 2).

As a final result of the relative orientation the projection rays aligned to the aforementioned four points should be on the same plane, then the volume of the tetrahedron they form becomes zero. This coplanar condition can be formulated by a determinant [1], [2].

$$\Delta = \begin{vmatrix} \bar{x}_{O_1} & \bar{y}_{O_1} & \bar{z}_{O_1} & 1 \\ \bar{x}_{O_2} & \bar{y}_{O_2} & \bar{z}_{O_2} & 1 \\ \bar{x}' & \bar{y}' & \bar{z}' & 1 \\ \bar{x}'' & \bar{y}'' & \bar{z}'' & 1 \end{vmatrix}$$

We remark that the exact equation contains a multiplier of $1/6$, but if we omitting this factor, it will not affect the result, since the volume should be zero in any way.

The image rotation angles $\omega', \varphi', \kappa', \varphi'', \kappa''$ as orientation elements are determining the orientation of planes and they form the rotation matrix elements for the images as direction cosines, which are built into the coplanar equation condition as trigonometric functions. So unfortunately, after the determinant is extracted, the resulting equation will not be a linear equation considering the orientation elements. The number of unknowns is 5, in order to solve the equalizations with the adjustment method of least squares there must be at least 6 stereoscopic points to be measured by the arrangement proposed by Gruber. Thus, according to Taylor series polynomial we solve the equations with a gradual approximation according to the method of least squares. We have only a few options to optimize the mathematical model. The only thing that could accelerate the iterative process is applying higher degree elements of the Taylor polynomial derivatives, although it is doubtful whether the greater computing need results more rapid iteration process.

The real change would be if we could describe an iteration-free process with the same number of unknowns. There are solutions by this approach mainly in the field of pattern recognition and robot vision based on artificial intelligence [3], [4]. In general, these solutions are derived with Gröbner-basis and there is no need to set up initial values for unknown parameters. This is, however, beyond the scope of the present paper.

C. Initial values

Since the conditional equations are generated by polynomial series (Taylor series), we need to know the initial values of unknown variables (orientation parameters). Normally it is not a problem to determine these initial values if the images have small rotation angles. In special cases, some terrestrial and aerial images can have large rotation angles, in which cases during the imaging we have to assure

and record the good approximations of the orientation elements.

III. CONCLUSIONS

As a summary we can say that the relative orientation is a necessary process but it would be better to avoid, since later we use only the collinear equations without any model coordinates. We have to go through the relative orientation

process only for calculation of γ''^0 as the theoretical vertical parallax.

Certainly, there are cases when we don't need geodetic coordinates and to know the model coordinates is sufficient. These tasks can occur in the terrestrial photogrammetry when we need the digital terrain model covering only a small area where the knowledge of scale factor is enough and we don't need to place the model into the geodetic coordinate system. The exact scale factor can be calculated by the reference lengths observable on images.

Another interesting field of research can be if we register the model coordinates although we don't need them for mapping but we could use them for gross-error detection.

REFERENCES

- [1] Albertz, J., Kreiling, W.: Photogrammetrisches Taschenbuch, Wichmann, 1989
- [2] Bronstejn, I.N., Szemangyejev K.A.: Matematikai szebkönyv, 3. kiadás, Műszaki Könyvkiadó, Budapest 1980.
- [3] Nistér D.: An efficient solution to the five-point relative pose problem, IEEE Transactions on Pattern Analysis and Machine Intelligence (PAMI), 26(6):756-770, June 2004
- [4] 70. Pan H.: A direct closed-form solution to general relative orientation of two stereo views, Academic Press, Digital Signal Processing, Volume 9, Number 3, July 1999, pp. 195-221(27), 1999

Interactive SQL Learning Course Software Specification

Zsobrák K., Gugolya L.

Obuda University/Alba Regia Technical Faculty, Székesfehérvár, Hungary

gugolya.laszlo@amk.uni-obuda.hu

zsobrak.krisztian@gmail.com

Abstract—In this paper a software's specification was developed, which is feasible in a later essay.

This software grants a possibility for more user groups, to learn the basic database management curriculum, in an interactive way with multiple SQL database engines (DBMS).

With this software the teacher can upload a specific module's tasks according to the themes, and the users will solve them. Unlike other solutions, it does not only compare the user's and the stored results, but it tests it multiple times, reducing the possibility of bypass the task.

KEYWORDS: MySQL, Oracle, SQL, interactive, software.

A. INTRODUCTION

In present days, the relevance of databases is people's everyday life is unquestionable.

Every website, which know a little bit more than just static information has to store its data in a database. Nowadays the most usual solution is MySQL [1-2], because it is freely accessible, and perfect to store smaller websites' and companies' data.

Therefor it is inevitable for a person working on the field of IT, to know the basics of database management.

In the age of information everyone can find the curriculum they need online, which they can use to learn these skills by themselves.

However the quality and the source of the curriculum is important. The university's education cannot be replaced by a website or a book.

This paper specifies a software, which can be used for individual learning, high school, or university level education.

B. MATERIAL AND METHOD

For the development of this study, we took a particular note of the existing SQL curriculums' structure, and interdependence [3].

The role of areas where this software is useful was important, besides the reduction of chance to bypass the task, which is achieved by testing the query on multiple (hidden) databases. This paper contains a study about managing simpler queries.

C. DEVELOPING THE CURRICULUM

To make this software, the occurrence of certain types of tasks has to be taken into consideration. The most usual

are the selecting, inserting, updating and deleting queries extended by conditions, groupings and joinings.

It is the teacher's task, to make the actual curriculum. He has to upload the proper databases, tasks, and solutions.

1) Existing, online solutions

Many websites are providing solutions like this, which can be found easily by someone who is willing to learn. These are based on existing structures, with static tasks. The goal of these SQL teaching websites is to introduce the user to the basics of SQL, rather than giving him permanent, real knowledge. It is obvious that the public teaching websites not capable of giving flexibility or verifiability, so they are less useful in education.

Some examples:

a) Codecademy:

This site has a clean [6], user-friendly interface, a good design, and also the elements are animated, as shown in a screenshot, in figure 1.

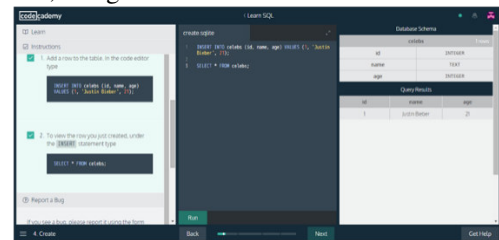


Figure 1. Codecademy

Its structure based on four main blocks:

Manipulation: Introduction to relational database management, creating and modifying tables.

Queries: Learning the most commonly used SQL commands to query a table in a database.

Aggregate functions: Learn how to use SQL to perform calculations during a query

Multiple tables: Learn how to query multiple tables using joins

In these blocks there are task categories, which contains multiple tasks that has to be solved in order to proceed. Next to the tasks there are descriptions of them, containing the conceptual background. First these are showing the solution, but later they expect the user to remember them.

The tests and the tasks are not changing, every user gets the same. The website does not provide features for creating personal tests, or verifying them.

b) SQLBolt:

The site is user-friendly [7], it reacts instantly and intelligently for the given queries, while also indicating if a table or field name is wrong, for example in figure 2.

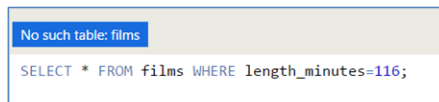


Figure 2. Example for wrong table name

Its structure expands from the basics queries, conditions, joinings, create tables, inserting and deleting.

It can only handle the syntax of MySQL. A screenshot is shown in figure 3.

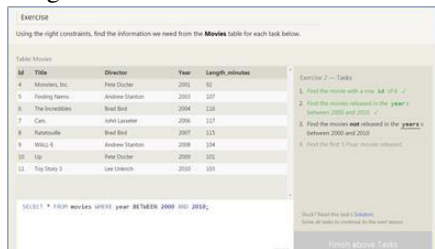


Figure 3. SQLBolt example

c) SQLzoo:

The interface of this website [8] is outdated, and simple structured, shown in figure 4.

It differs from the others because here the user can select the database management engine, which he wants to use, but the curriculum is much smaller than the previously mentioned websites'.

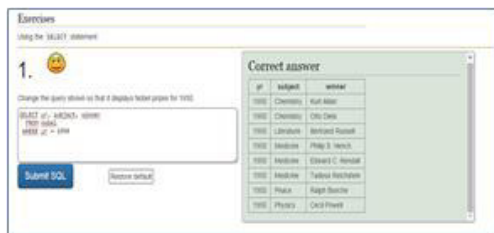


Figure 4. SQLzoo example

2) Other solutions, comments

There are existing solutions for making online tests (for example: Moodle), but these are not supporting the checking of SQL queries. They only compare the user's solution with the existing one letter by letter.

The SQL teaching site, specified in this study is similar to the existing solutions, however it would expand their functions.

The teaching process starts with the teacher, who defines which tasks he would like the students to solve. Because the teacher is the one, who uploads the tasks, the database and the solutions, the flexibility and difficulty of the tasks are highly increase.

The instructor can monitor how the students succeeded on a specific task. He can check the progression, which modules the students have been succeeded, and which requires more explanation. This statistics can help the teacher giving grades.

Because the curriculum depends entirely on the teacher, he can create themes, difficulty or elements that are required, without external help.

The existing modules can be expanded later on, which means there is no need to re-upload everything.

D. DISCUSSING THE PROBLEM

1) Applications

In order to use the software dynamically, multiple groups of users must be targeted.

We defined three areas, where this software could be used:

a) For external user

This level is equal to the mentioned concurrent technologies. A user registers to the system, then he is able to solve tasks in the modules. This solution could be used by the institution to make placement tests, academic competition, or simply raise the institution's reputation.

b) High school level

This system could be used to upload mid-level or advanced level graduation tests. The students can prepare by practicing on previous years' exercises, and the teacher can see, which modules need more study. It also can be used on lessons to make tests or exams.

c) University level

In higher education the database management courses based on high school's curriculum, but it's not required from the students to have this knowledge. In the university, the first half of database management courses are the same as the high school courses, but later in the second half new technologies are introduced, like triggers and stored procedures. This system can help to students to practice and prepare for lessons or write exams.

2) Checking solutions

To check the user's solution is a key problem in the software. If a teacher is checking a solution, he takes into consideration, that how the student was thinking and did he understood the key connections is the task.

If a computer checks a solution, this method would be exceptionally hard. Creating such software would be difficult, time and resource consuming. In these environment the most often used method is the comparing of the results for the queries, which can be done easily by a computer. For this, the system needs to know, what the correct answers are.

The result of the teacher's query can be stored, which are reducing the computing needs. But this method requires storage, and would be a less flexible solution.

It is more effective, if the solution's query is stored, which has to be given by the instructor. In this case both the instructor's and the user's query must be processed, and the results are compared. If the difference is zero, the solution is correct. It requires that fields are given in the same order, or the field names to be matching.

3) Resolve language differences

Nowadays, the most common database management system in the business sphere is Oracle, which provides a lot of useful features. The second most common DBMS is MySQL, which is functionality limited, but greatly

satisfies the needs of usual web-based services, and it is entirely free. This is the main reason why most of the websites use MySQL for database engine.

Therefore, these systems are worth using to create tests, because the user will see these with the highest possibility.

If we give the opportunity to the user to select which database engine he would like to use for solving the tests, then checking the query is a harder task.

The syntax of Oracle and MySQL queries are basically the same, but some functions can be accessed in a different way.

There is a chance for a perfect MySQL query to generate error in Oracle. There are some examples to resolve this problem:

a) *Converting queries*

There are solutions to convert queries from Oracle to MySQL syntax, vice versa [4]. These softwares are first interpreting the query and transform it by its knowledge of differences. These solutions are usually paid desktop applications, which are difficult to integrate to a web-based system.

b) *Storing multiple solutions*

Another solution for this problem is to store the teacher's query in multiple syntax. In this case when he enters a task, he declares the solution in every accessible syntax. When the user logs in, he must be able to choose which database system he wants to use, and the software will run his queries on that.

4) *Reduction of cheating*

Because the system evaluates the users work based on already solved tasks, some methods has to be defined in order to prevent the possibility of cheating.

a) *Copying*

Every user gets the same tasks, and this also gives them a chance to cheat. For example, if a user gets a solution and turns it in as his own. A filter can be developed, which checks the similarity of the queries. But results of the filter must be handled with care, because the solutions often have similar structures, resulting two users can write the exactly same query without knowing each other.

Another solution is to make a program, which listens the pressed keys in the browser. This can indicate that the user just copied the query in with some shortcuts, or wrote it there. This method can give false results, if the user's browser cannot run this program or the user made the query in a different editor, and copied it into the webpage.

b) *Bypass*

There is also a chance of cheating, if the user knows the data which are stored in the database's tables, and bypass the task by writing a query, which results in that his solution will be the same as the teacher's, without the knowledge that the task requires. For example, if a task is about filtering data with multiple conditions, but the user simply filters for the identification number of the rows.

In this example from SQLBolt, the user has to find every Toy Story titles. The system introduces the using of the "LIKE" keyword, but this task can be easily bypassed with a trick.

The user first makes a query which lists all data from the table. He searches for the expected rows, and notes the identification numbers (id). He filters for the id, and the system accepts the answer, as shown in figure 5.

id	Title	Director	Year	Length, minutes
1	Toy Story	John Lasseter	1995	81
3	Toy Story 2	John Lasseter	1999	93
11	Toy Story 3	Lee Unkrich	2010	103

Exercise 3 -- Tasks

1. Find all the Toy Story titles ✓
2. Find all the titles directed by John Lasseter
3. Find all the titles not directed by John Lasseter
4. Find all the WALL-E movies

SELECT * FROM movies WHERE id IN (1,3,11);

Figure 5. SQL ID Fraud

A solution for this problem is that we create private databases, which the user can't access, or does not know about, so he cannot bypass the task.

There are two methods for creating a private database:

c) *Reduced database*

The teacher uploads the whole database when he uploads the modules, which have to be properly large. Then he sets up the connection between the tables. This will be the private database.

After that, the software reduces its size, while paying attention to keep the records which are linked by foreign keys. This reduced size database will be the public one, which the user can write his queries.

d) *Random database*

Another solution is that the system generates the private database from random data. This method also requires the teacher to upload the database and set up the connections between the tables. After that, the system analyzes it, and creates statistics from the data stored in it. This is the base of the random generation. The random database has the same structure, table names, field names and types, only the data is random. This isn't required to be sensible, because neither the users nor the teacher will see it.

Because only the query is stored as a solution, the question is, that the results of the teacher's and the user's query are the same.

5) *Analyzation of the query's performance*

Comparing two queries, the execution times, the used memory, etc. can also be taken into account. If the user's query satisfies all of the requirements in all aspects, but its performance is significantly worse, that indicates that he used a technique wrong.

In this case, the teacher can deduct some points.

But this cannot be easily analyzed, because the database management system is basically made up for optimization. If a query runs for the first time, it requires more resources, than the second or third run, because the results are stored in a cache, accelerating the execution.

It is highly possible, that the user's and the teacher's query do the same thing, and the examination of differences will not give relevant information.

E. CONCLUSION

The software discussed in this paper does not exist yet, but there would be a claim for it. Nowadays the students solve the tasks and tests on real database engines, but checking the solution is done manually by the teacher. The

existing e-learning systems can provide text input or multiple choices, but these are not capable of checking the SQL query. However the publicly accessible teaching websites cannot be utilized in education.

A software like this could help the teachers, because it makes easier to check the students' solutions, and it would require less time and energy.

It can help to high school students to prepare for graduation. Although Microsoft Access [5] gives opportunity to use design view to create queries, which can be done without using SQL, but this solution works only for Access. If a person wants to use another database management system, the knowing of SQL is key.

To introduce the software in education, the teachers must be prepared to use it, uploading tasks, solutions, and how the results can be evaluated. There wouldn't be a big

difference in the lessons, because there are already e-learning system in educations. But it would ease the teachers' work, and the lessons would be more dynamic.

REFERENCES

- [1] <https://dev.mysql.com/doc/>, (19.08.2015).
- [2] <https://www.oracle.com/solutions/index.html>, (19.08.2015).
- [3] <http://db-engines.com/en/>, (18.08.2015).
- [4] <http://www.sqlines.com/home>, (29.08.2015).
- [5] <https://products.office.com/hu-hu/access>, (19.08.2015).
- [6] <https://www.codecademy.com/learn/learn-sql>, (23.08.2015).
- [7] http://sqlbolt.com/lesson/select_queries_introduction, (23.08.2015).
- [8] http://sqlzoo.net/wiki/SQL_Tutorial, (24.08.2015).

Quality Assessment System in Obuda University Alba Regia Technical Faculty

Margit Horosz-Gulyás, PhD

Alba Regia Technical Faculty of Óbuda University, Székesfehérvár, Hungary
horoszne.margit@amk.uni-obuda.hu

Abstract—What is important for the employee in its workplace (especially in higher education)? Why should the management know what motivates the staff? What is the role of the students? These are only a few questions, I looking for the answers. Due to a project a new assessment system was made based on the motivation of the employees.

Keywords: Quality, Employee, Motivation, Students

I. INTRODUCTION

We work in surroundings that our colleagues of thirty years ago would not recognize. Higher education has become part of a global shift to a new way of creating and using knowledge. The new way is focused on solving problems and be sensitive to customer needs. It strives for quantity as well as quality. It cuts across disciplinary boundaries.

In knowledge-based economies, governments see universities as engines for social change and the expansion of prosperity. University teachers have accordingly found themselves working harder and at the same time being required to be more businesslike and more accountable. The pleasures of the academic life have dwindled for many university teachers. They are unimpressed especially by the administrative effort associated with quality assurance and accountability. It uses up time and energy that could be focused on the core business of research and teaching [1].

Practically everybody in the academic community gets assessed these days, and practically everybody assesses somebody else. Students, of course, come in for a heavy dose of assessment, first from admissions offices, later from the professors who teach their classes, and increasingly from administrators complying with state accountability requirements. Students are also active participants in the assessment business, with end-of-course evaluations that are widely used by colleges and universities and various forms of web-based assessments of professors. The whole institution is regularly assessed in a highly detailed fashion by external accrediting teams made up of faculty and administrators from other institutions.

As commonly used today, the term assessment can refer to two different activities: the gathering of information (measurement) and the use of that information for institutional and individual improvement (evaluation) [2].

To manage change, universities and also other organizations must have employees committed to the demand of rapid change and as such committed employees are the source of competitive advantage. The concept of human capital and knowledge management is

that people possess skills, experience and knowledge, and therefore have economic value to organizations. These skills, knowledge and experience represent capital because they enhance productivity. Motivations represent those psychological process that cause the arousal, direction and persistence of voluntary actions that are goal oriented. There are several motivation theories. Maslow's need theory was the development of the hierarchy of needs. He believed that there are at least five sets of goals which can be referred to as basic needs and are physiological, safety, love, esteem and self-actualization. Maslow stated that people, including employees at organizations, are motivated by the desire to achieve or maintain the various conditions upon which these basic satisfactions rest and by certain more intellectual desires.

One of the earliest researchers in the area of job redesign as it affected motivation was Frederick Herzberg. He and his associates began their initial work on factors affecting work motivation in the mid 1950's. Herzberg discovered that employees tended to describe satisfying experiences in terms of factors that were intrinsic to the content of the job itself. These factors were called motivators. Motivators (e.g. challenging work, recognition for one's achievement, responsibility, opportunity to do something meaningful, involvement in decision making, sense of importance to an organization) that give positive satisfaction, arising from intrinsic conditions of the job itself, such as recognition, achievement, or personal growth.

TABLE I.
HERZBERG-FACTORS

Motivators	Hygiene factors
satisfaction with the work	salary
challenge by the work	coworker relations
good team	safety
enough information	authority
	enough information
	work environment
	training, career
	spending leisure time
	solve problems, complaints

Conversely, dissatisfying experiences, called hygiene factors, largely resulted from extrinsic, non-job related factors, such as company policies, salary, coworker relations, and supervisory styles [3]. Hygiene factors (e.g.

status, job security, salary, fringe benefits, work conditions, good pay, paid insurance, vacations) that do not give positive satisfaction or lead to higher motivation, though dissatisfaction results from their absence. The term "hygiene" is used in the sense that these are maintenance factors. These are extrinsic to the work itself, and include aspects such as company policies, supervisory practices, or wages/salary.

II. THE RESEARCH

The aim of the research was to make a quality assessment system to Alba Regia Technical Faculty. In order to make that we create a questionnaire in which the motivations of the employee was estimated. The questionnaire was divided up to six parts.

In the first part there were personal data (birthdate, workplace-institute etc.). The second part was general, questions about the satisfaction with the workplace, the environment of the workplace etc. In the third part we assessed the motivations of the employees using Herzberg factors: intrinsic and extrinsic factors (motivators and hygiene factors). In the next part of the questionnaire we tried to estimate the talent of the employees. The aim of the next part was to measure the satisfaction with the teaching. The last part referred to the estimate the satisfaction with the assessment system.

24 colleagues filled the questionnaire, 60% was female, 40% male. 13% was leader, 87% was employee. The distribution of the birth can be seen in Table 2.

TABLE II.
DISTRIBUTION BY THE BIRTH

Birth date	Person	
1945-1964	11	46%
1965-1979	10	42%
1980-1994	3	12%

The colleagues ranked what is important for their on a workplace out of the next factors: salary, work environment, training, career, bonus (Fig. 1.)

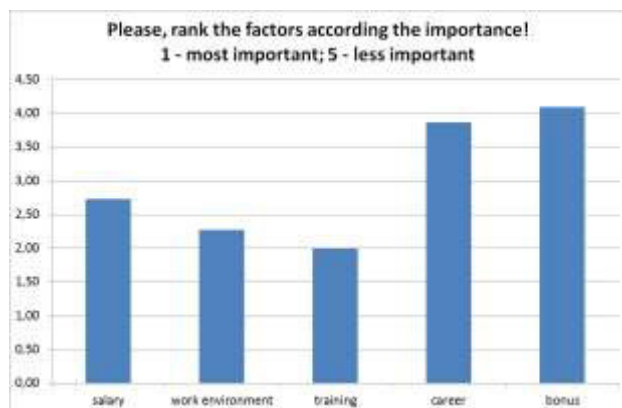


Figure 1. What is the most important for you in your workplace?

The survey was verified that for the colleagues is very important the training. The next factor was work environment: the most colleagues like to work alone (as teacher in general), and in a separate room. These things

are commonly characteristic the researchers. The third one was the salary, but it is true that as public servant it is not a motivator.

We got the same result in other analyses, in which we research what is the most important strong point in the workplace (Fig. 2.) The most important factor were the stability, the balance between the work and private life and the respect the colleagues. These factors are in essence the work environment. The less important factors were the outing and creature comforts.

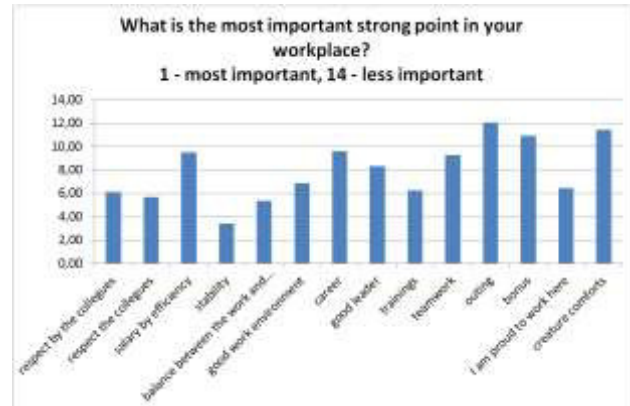


Figure 2. What is the most important strong point in your workplace?

To survey the motivation area among the colleagues were used the Herzberg factors. The questions were: How satisfied are you if..., and How dissatisfied are you if... (Table III, IV.)

TABLE III.
HYGIENE FACTORS

control	methods of the control
work specifications	technical specifications work environments
coworker relations	coworker relations relations with the leaders
salary and safety	work safety salary, bonus work-life balance
company policies	justice, legality

TABLE IV.
MOTIVATORS

achievement	deliverance
appreciation	direct feedback leader appreciation
work	work flow work diversity great importance work
responsibility	responsible work
career and growth	career

III. THE RESULTS

The first step was to analyses how dissatisfied the colleagues are if any factor is missing. In Table V. in the first column we can see the ranked statements (in Likert-scale 1-5). In the second column there are the average values of the statements, the next are the deviations (the average of the deviation is relatively high, 1,22); and in

the last column there are the sign of the Herzberg-factors (M-motivator, H-hygiene factor).

The survey verified us the hygiene factors, like safety, salary by efficiency, balance between the work and private life, the relationships with coworkers and leaders, work environment. The unfitting factor was achievement by the work, which can be due to the technologies and modern work logistic. Everybody work a lot and achieve in high level, so it is naturally if somebody good at work. The other factors were responsibility and great importance of work which are bound up with the achievement.

TABLE V.
DISSATISFACTION

How dissatisfied you if...	Average	Deviations	Factor
the work is uncertain	4,17	1,18	H
the salary is not by efficiency	4,08	1,11	H
not enough private life	4,04	1,06	H
not the best achievement	3,96	1,17	M
the coworker relationships are not good	3,88	1,33	H
not good relationship with the leader	3,71	1,51	H
not good work environment	3,63	0,99	H
you do not feel the importance of the work	3,63	1,32	M
not enough responsibility	3,50	1,35	M
not good company policies	3,42	1,15	H
you do not get the leader appreciation	3,38	1,18	M
no deliverance	3,38	1,28	M
not good control	3,25	1,33	H
the work is not diverse	3,25	1,20	M
no feedback	3,21	1,26	M
no career	2,96	0,93	M
the output of the work is hard to identify	2,96	1,34	M
lack of importance of the work	2,21	1,29	M
	3,48	1,22	

The lack of motivators does not cause dissatisfaction (below in Table V.) To control this statement we made another table (Table VI), in which we can find the answer to the question: How satisfied are you if....

We can wait that the above mentioned (Table V.) motivator factors will be in the Table VI. The first result is that the good achievement is the best important for the colleague. The other group of the motivators (career, feedback) has no high score both table, so these factors neither are motivators or hygiene factors, the employee do not calculate with these factors.

The method of the control and salary are hygiene factors. Motivators are responsibility and responsive work. Other motivators are leader appreciation and the good identified outputs.

The relationships with coworkers and leaders are motivators and hygiene factors also. That means, that

someone is missing, that can cause dissatisfaction, and if anything is very good, that can cause satisfaction. Also the balance with private life, the safety and work environment are motivators and hygiene factors also. Individuals spend a lot of time at work. It is safe to say that work may seem like their home away from home. Since they spend so much time with coworkers and management, they need a pleasant work environment.

TABLE VI.
SATISFACTION

How satisfied you if...	Average	Deviations	Factor
good achievement	4,67	0,75	M
good relationship with coworkers	4,54	0,76	H/M
good relationship with the leaders	4,46	1,00	H/M
safe work	4,42	1,04	H/M
the output can be easily identify	4,33	0,80	M
the work is diverse	4,25	0,78	M
good work environment	4,04	0,98	H/M
good company policies	4,04	0,93	H/M
enough private life	4,00	1,22	H/M
enough leader appreciation	3,96	1,10	M
good deliverance	3,96	1,21	M
responsible work	3,92	1,19	M
responsibility	3,92	1,08	M
salary by efficiency	3,92	1,22	H
good control methods	3,75	1,09	H
direct feedback	3,58	1,04	M
career	3,13	1,27	M
enough importance of the work	2,46	1,35	H
	3,96	1,04	

IV. CONCLUSION

Workplace satisfaction helps companies extract the best performance from their employees. According to Herzberg, hygiene factors are what cause dissatisfaction among employees in a workplace. In order to remove dissatisfaction in a work environment, these hygiene factors must be eliminated. There are several ways that this can be done but some of the most important ways to decrease dissatisfaction would be to pay reasonable wages, ensure employees job security, and to create a positive culture in the workplace. Eliminating dissatisfaction is only one half of the task of the two factor theory. The other half would be to increase satisfaction in the workplace. This can be done by improving on motivating factors. Motivation factors are needed to motivate an employee to higher performance.

The survey shows that the negative events have more influences to our satisfaction/dissatisfaction. That shows

the questionnaire and the tables above. Some suggestions are:

- Creating complete and natural work units where it is possible. An example would be allowing employees to create a whole unit or section instead of only allowing them to create part of it.
- Providing regular and continuous feedback on productivity and job performance directly to employees instead of through supervisors.
- Encouraging employees to take on new and challenging tasks and becoming experts at a task.
- Not receiving feedback on their work can be quite discouraging for most people. Effective feedback will help team members know where they are and how they can improve.
- Autonomy and control are necessary for people to feel satisfied with their work. In fact, psychologists have found that the less control people have over their jobs, the more stressful and unsatisfying they find it.

ACKNOWLEDGMENT

We acknowledge the financial support of this work by the Hungarian State and the European Community under

the TÁMOP-4.2.2.B-15/1/KONV-2015-0010 project entitled "Development of Alba Regia Technical Faculty's Scientific Workshops".

REFERENCES

- [1] P. Ramsden: Learning to teach in higher education. Routledge, 2003, 265 p. https://books.google.hu/books?hl=hu&lr=&id=As-IK_4wAJz4C&oi=fnd&pg=PR1&dq=quality+assessment+method+higher+education&ots=XKXGRNkaK1&sig=JWWaxQngFHFg2phSbCcQyE5RrH0&redir_esc=y#v=onepage&q&f=false
- [2] A.W.Austin, A. L. Antonio: Assessment for excellence: The philosophy and practice of assessment and evaluation in higher education. Rowman and Littlefield, 2012, 367 p. https://books.google.hu/books?hl=hu&lr=&id=jEsRVtbukyMC&oi=fnd&pg=PR5&dq=benchmark+higher+education&ots=DJUbTSH33j&sig=vh5i5jNgxJkvt4KVk7imGyvqFKQ&redir_esc=y#v=onepage&q=benchmark%20higher%20education&f=false
- [3] S. Ramlall: A review of employee motivation theories and their implications for employee retention within organizations. Journal of American Academy of Business, Cambridge, Sep 2004, 5, ½. pp. 52. – 63. ftp://118.139.161.3/pub/moodledata/113/Ramlall_2004.pdf

Measures to Prevent the Impact of Carcinogenic Asbestos during Demolishing of Buildings

Gencho Panicharov*, Aneta Georgieva*, Marta Seebauer**

*Varna Free University "Chernorizets Hrabar"/ Faculty of Architecture, Varna, Bulgaria

**Óbuda University, Alba Regia Technical Faculty, Székesfehérvár, Hungary

dr_panicharov@abv.bg, angr@abv.bg, seebauer.marta@amk.uni-obuda.hu

Abstract — The article reviews the risk factors of the effect of the asbestos on the health of the workers. The strongly carcinogenic effect of asbestos has been highlighted. Organizational measures to protect workers during demolition of buildings have been proposed. The necessity to keep the population safe during transportation and storage of materials containing asbestos is strongly emphasized.

I. INTRODUCTION

Asbestos belongs to a group of ten substances established to have the greatest carcinogenic potential for malignant cancers such as two types of lung cancer (lung tissue and mesothelioma), pharynx cancer, larynx cancer, colon cancer and stomach cancer [1]. It is listed as the third most abundant pollutant in the global scale [2].

The health risk associated with exposure to asbestos depends on the type of mineral used in the product, concentration of asbestos dust, size of inhaled fibres and time of exposure. From many factors enhancing negative effects of asbestos fibre presence in ambient air the most important is cigarette-smoking.

The possibility of exposure to commonly applied asbestos and other mineral fibres has led to the introduction of legal regulations significantly restricting or completely forbidding their use. Also, the use of materials containing asbestos is strictly controlled. In world, asbestos has been listed in a group of 5 substances whose use is heavily fined. Much effort has been made to replace asbestos by environmental-friendly materials and to devise effective methods of its elimination.

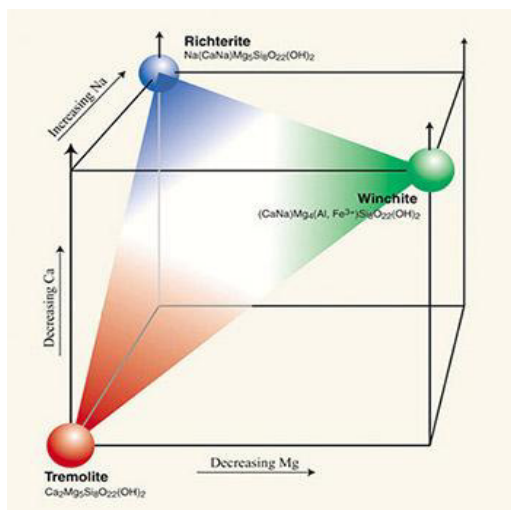


Figure 1. Components of asbestos



Figure 2. The electron microscope image of chrysotile asbestos $3\text{MgO} \cdot 2\text{SiO}_2 \cdot 2\text{H}_2\text{O}$

II. CHEMICAL AND PHYSICAL CHARACTERIZATION OF ASBESTOS AND YOUR NATURAL FIBRES

According to the chemical composition of asbestos minerals, they can be classified as hydrated silicates of magnesium, iron, calcium, sodium and other elements [3].

These minerals differ by the content of the main components (SiO , FeO , MgO , CaO) – Fig.1, colour (from white – Fig. 2; to pale green, yellow, pink, blue – Fig. 3, pale grey, pale brown, dark green – Fig. 4), density, melting point (from 1200 to 1500°C) resistance to acids and bases, arrangement of fibres (i.e. texture - elastic, silk-like and hard) and mechanical properties (resistance to stretching and rotation).



Figure 3. The electron microscope photo of crocidolite asbestos $(\text{Na}_2\text{Fe}_3^{2+}\text{Fe}_2^{3+})\text{Si}_8\text{O}_{22}(\text{OH})_2$



Figure 4. The electron microscope photo of amosite asbestos ($\text{Fe}^{2+}, \text{Mg})_7\text{Si}_8\text{O}_{22}(\text{OH})_2$)

Commercially available asbestos minerals include: - chrysotile (white asbestos) one of the minerals most commonly met in the earth crust; its melting point 1500°C , poor conductor of heat, electricity and sound; resistant to high temperatures because of a high concentration of magnesium oxide. This mineral is characterized by uniform chemical composition the metal-oxygen layer made of the oxygen octahedrons with different cations (mostly magnesium) inside – Fig. 5. These packets are neutral macromolecules, which are arranged into the crystal structure.

Asbestos, meaning indestructible or unquenched in Greek, is the commercial name applied to a group of crystalline and relatively insoluble silicates of chain or ribbon structure. The group includes 6 minerals one of which occurs in the silicates of the serpentine group characterized by the chain structure and the others belong to amphiboles with ribbon anions.

A representative of the first group is chrysotile, and the other includes crocidolite, amosite, antophyllite, tremolite and actinolite.

The bonds within the silica-oxygen chains in the two above-mentioned silicate structures are relatively strong.

Asbestos, which was already known in prehistoric times (1800 BC), is a material which never burns, rots nor corrodes. It insulates from cold and noise, it possesses both high elasticity and tensile strength. In short, asbestos materials are virtually indestructible. Because of its' excellent characteristics, the mineral asbestos has been used in countless ways in technologies of every kind. As a rule, a distinction is drawn between weakly bound asbestos products such as: asbestos sprayed materials, asbestos insulating board, asbestos fabric; and strongly bound asbestos products such as asbestos-cement sheets, corrugated asbestos-cement sheets, asbestos-cement pressure pipes.

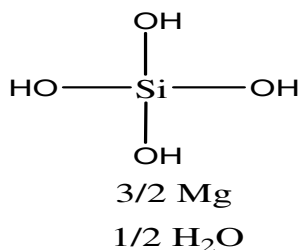


Figure 5. Organic structure of asbestos minerals

III. IDENTIFICATION AND QUANTIFICATION OF MINERAL FIBRES IN AMBIENT AIR

The presence of asbestos minerals in ambient air can be detected via pertinent colour tests dyeing asbestos fibres in shades of blue or red. In contrast to the majority of harmful substances in ambient air, the presence of inorganic fibres can be easily determined (abundance and kind) by optical and electron microscopy observing the deposits on membrane filters after air filtration.

The methods for determination of asbestos fibres differ in the mode of sample collection, instruments used and techniques of fibre counting. The fibres deposited on membrane filters can be performed by the following methods:

- using a phase-contrast optical microscope (PCM), for the first time applied in 1966, assuming as dangerous the fibres of the length to diameter ratio $> 3:1$, which corresponds to a length greater than $5\text{ }\mu\text{m}$ and a diameter smaller than $3\text{ }\mu\text{m}$; the fibres represent the biologically active part of respirable fraction (penetrating the air sacks);
- transmission electron microscope (TEM), equipped with an energy dispersive X-ray analyzer;
- X-ray diffractometer (XRD);
- scanning electron microscope, coupled with an energy dispersive X-ray diffractometer (EDXA);
- selected area electron diffraction (SAED);
- laser microprobe mass spectrometry (LMMS);
- polarized light microscopy (PLM).

IV. MECHANISM OF FIBRE CARCINOGENESIS

The first reports on the development of lung cancer as a result of inhalation of asbestos-containing dust prompted intense studies aimed at an explanation of the mechanism of carcinogenic effect of asbestos fibres. Although it has been established that the potential of the carcinogenic effect of the fibres depends on their chemical stability and size, not one of the hitherto proposed mechanisms has not been commonly accepted (see Fig. 6).

The reason is the lack of convincing experimental evidence from cell tests which would give a direct connection between the initiation of cancer changes and the interaction of asbestos fibres with DNA.

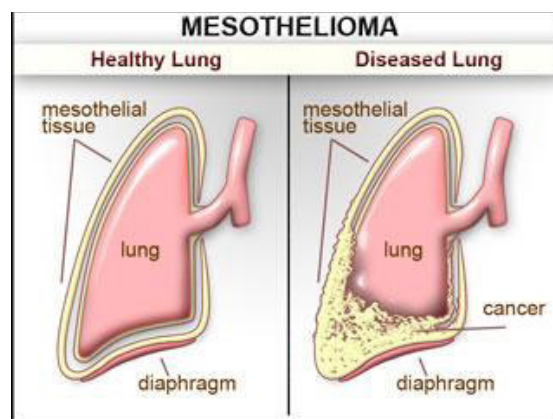


Figure 6. Aisease mesothelioma caused by asbestos

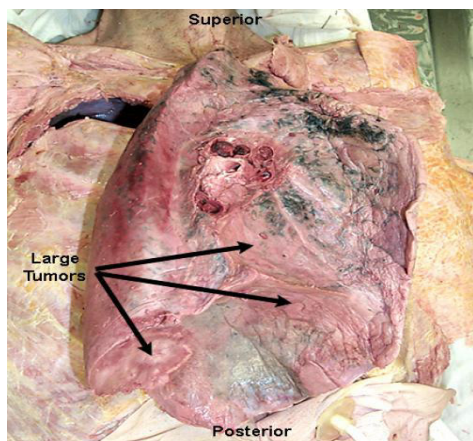


Figure 7. Tumors caused by asbestos fibres

One of the proposed mechanisms described as the hypothesis of a carrier explains the appearance of cancer changes by the fact that asbestos fibres can adsorb and carry polycyclic aromatic hydrocarbons (PAH) to the lungs. The fibres of asbestos can also interact with the proteins of red blood cell membrane leading to dissolution of red blood cells (haemolysis) (see Fig. 7).

Results of many studies prove that heavy metals substituted in crystalline structures of asbestos, mainly iron, and PAH adsorption are particularly dangerous for man when an air sack in his lungs has been pricked by microfibrils of asbestos carrying PAHs or metal ions.

The hemolytic activity of asbestos is determined by the type of the mineral, the ratio Mg:Si and the size of the fibre area. The greatest biological activity is believed to characterize chrysotile fibres. There is increasing evidence that iron from asbestos may cause cancer as a result of radical induced damage. Many mechanisms have been proposed as the potential reactions of iron to catalyze the oxidation of DNA, lipid and protein. Analysis suggests that different chemical reactions catalyzed by asbestos and other mineral fibres, e.g. erionite, are responsible for their pathological influence.

V. METHODS OF ELIMINATION OF MINERAL ASBESTOS AND THEIR MODIFICATION INTO HARMLESS PRODUCTS

The first measures taken to eliminate or limit exposure to asbestos consisted in wet separation and ventilation of the rooms containing sources of its emission, appropriate organization of such activities as removal, transportation and storage of asbestos waste and conservation works.

The simplest measures aimed at elimination or limitation of emission of asbestos dust were mechanization of all kinds of handling of asbestos-containing materials, which should be wet processes if possible.

Emission of asbestos dust can essentially be limited if the materials containing asbestos are wetted with surfactants whose use also limits the use of water. For example, an anionic surfactant, such as linear alkyl sulphonate, sodium lauryl sulphates, polyalkoxy carboxylates, or an non-anionic surfactants, such as alcohol alkoxylates, alkyl phenol ethoxylates, polyoxoethylene esters and polyalkene oxide block

copolymers can be added to the treating composition in conventional amounts, e.g. up to about 5% by weight.

The problems of wet removal techniques include the requirements of physical removal and handling of the wet asbestos-containing material. In addition, the removed material must be further treated to destroy the remaining asbestos component if the material is to be discarded as a non-asbestos containing material.

VI. ANALYSIS TECHNICAL PROPERTIES OF ASBESTOS

Asbestos materials are very widespread around the world as it is shown in Table 1.

Asbestos has extraordinary valuable technical properties - very high thermal insulation characteristics, fire resistance, mechanical strength and others. That is why asbestos is widespread and has been found very useful in construction.

Exceptional technical properties of asbestos are the following:

- high tensile strength;
- high length / diameter (aspect) ratios minimum of 20 – can be up to 1000;
- sufficiently flexible to be spun; it is more suitable for weaving into a material and has been preferentially used in textiles and paper products;
- effect of high temperatures on tensile strength - largely unaffected up to 550 °C, above this temperature dehydroxylation occurs with resulting rapid loss of tensile strength;
- combustibility - asbestos fibres do not burn, although will undergo changes at high temperatures; widespread use as fire-proofing
- thermal conductivity - all asbestos types have very low thermal conductivities – i.e. they are all very good insulating materials; widespread use in thermal insulation and lagging;

Research in medicine has shown that the asbestos fibers cause cancer and other diseases. That is why asbestos is a prohibited mineral in the industry and construction.

TABLE I.
1995 WORLD PRODUCTION (METRIC TONES)

Commonwealth of Independent States	1,000,000
Canada	510,800
China	250,000
Brazil	180,000
Zimbabwe	145,000
South Africa	100,000
Greece	50,000
Swaziland	30,000
India	25,000
United States	9,000
Colombia	5,000

Asbestos fibers associated with these health risks are too small to be seen with the naked eye. Breathing asbestos fibers can cause a buildup of scar-like tissue in the lungs called asbestosis and result in loss of lung function that often progresses to disability and death. Asbestos also causes cancer of the lung and other diseases such as mesothelioma of the pleura which is a fatal malignant tumor of the membrane lining the cavity of the lung or stomach. Epidemiologic evidence has increasingly shown that all asbestos fiber types, including the most commonly used form of asbestos, chrysotile, causes mesothelioma in humans [4].

CONCLUSIONS

The following question has been raised: should this very valuable mineral be forgotten from the technical specialists.

The answer is no. Methods and ways for disposal of asbestos fibers should be researched and explored.

Then this very valuable material will find widespread use in industry.

This way we will thank nature for granting us with the existence of a mineral that has such extraordinary valuable technical properties.

REFERENCES

- [1] Asbestos Fibers and Other Elongate Mineral Particles: State of the Science and Roadmap for Research DHHS (NIOSH) Publication Number 2011-159 .March 2011
- [2] Herbert R, Moline J, Skloot G, et al. The World Trade Center disaster and the health of workers: five-year assessment of a unique medical screening program. Environmental Health Perspectives 2006
- [3] Skammeritz, E "Asbestos Exposure and Survival in Malignant Mesothelioma: A Description of 122 Consecutive Cases at an Occupational Clinic." The International Journal of Occupational and Environmental Medicine (IJOEM), Vol 2, No 4 October 2011.
- [4] O'Reilly KMA, McLaughlin AM, Beckett WS, et al. Asbestos-related lung disease.American Family Physician 2007

Adaptive Version Control in ERP Environments

A. Selmeçi

Óbudai Egyetem AMK, Székesfehérvár, Hungary

E-mail: selmeçi.attila@amk.uni-obuda.hu

Abstract—The modification of anything brings changes in a system. If we consider only the IT world we can think about hardware and software elements, components or even huger systems built from other elements or components. An Enterprise Resource Planning system, which is the central heart of a company today, is not a single, simple solution, but a very complex, integrated environment. The integration takes part internally, like database system or web servers; externally like interface based communications with other solutions; and technically like using storage, network, hardware, etc.

The modifications on different level build versions and these versions should be applied occasionally or in a regular manner. The goal in this paper to show the possibility of hardware-software solution, which is adaptive enough to obscure the modifications brought by versions. In the same time we introduce the availability of analogous versioning as well.

Keywords: ERP, version management, multi-tier, storage, virtualization, redundancy

I. INTRODUCTION

The software applications can be grouped according to the size, functionality, architecture, etc., but almost all has software implementation bugs. These bugs should be corrected some times and the software developers provide new releases, versions, or even bug fixes, patches for the given software.

In the real life the machines (simple or complex) have also components, mounting parts, which should be repaired or sometimes replaced by a new one. We can think about corrosion, physical damages, using of pure material, inappropriate design and manufacturing, etc. Many cars, or other machines are recalled on a designing bug, which comes out only in the special cases occurring in daily usage.

The hardware elements around the software applications can be regarded as such special products, like cars. These can have issues, problematic building elements or any component, which could fail.

The producer, manufacturer repairs the bug, or issue, or even problem within the original software application, machine, or hardware. They should internally test it and provide the solution in a patch, or in a new release. Sometimes the new releases are not really implementable, but the new one should step pin the place of the original one. In software solutions would we can ask for a migration path between the two versions, but in hardware environments or among used machines total cutover is needed as well.

The next chapters of this paper describe the versioning types, occurrences in business application environment and the possible theoretical future and implementation of a sustainable, well-managed version control.

II. ERP SYSTEMS AND THE VERSION MANAGEMENT

In case of a business application, like Enterprise Resource Planning (ERP) systems, we have more complex environment and co-existing sources of errors. An ERP system has nowadays at least three-tier architecture. [4] It contains a database engine, an application layer, where the business logic runs and a frontend layer. In case of web application there could be more tiers. [5] Each tier has its own purpose and internal architecture. Fig. 1 visualizes a standard three-tier ERP and a four-tier web enabled ERP architecture.

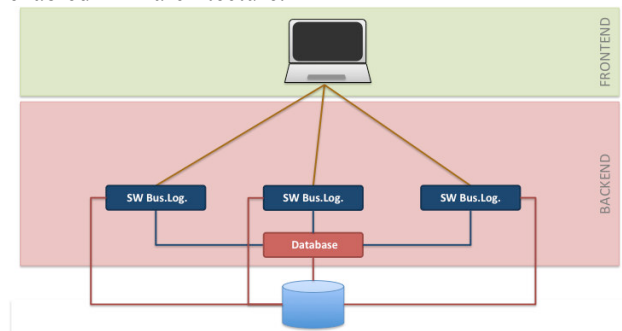


Figure 1 – ERP simplified architectures

These components of the multi tier architecture can be application specific or general-purpose applications. Let us consider for example a web-based software application or ERP solution having three or four tiers. To store the data many ERP systems use a database management system (DBMS). Some installations are database dependent, like applications using open source technologies (e.g. MySQL). Others are database independent or rather they can use, support not only one database management system, but several ones (like SAP does with Oracle, MSSQL, DB2, Informix, MaxDB, Sybase ASE or even HANA DB). For the web enablement there should be somewhere a web server, which provides the web contents, stores the mime objects and handles the HTTP requests. The same is true for this web layer as for the database layer. There are some special ERP applications, which contain their own internal web server implementation, but others use one from the market (e.g. Microsoft Internet Information Server, Apache HTTP Server). These standard, public solutions (DBMS and web server) are separately developed, tested, distributed and supported. They have their own version management and version control inside.

The Fig. 2 below sketches the main software and hardware layers, component under an ERP system.

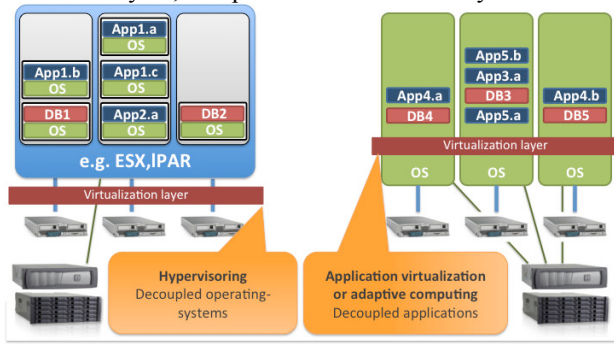


Figure 2 – Undelaying components supporting ERP

Each of the applications should run on a computer or node, which can be physical or virtual as well. There are standard open source operating systems and proprietary ones. The non open-source operating systems mainly have their own hardware as well (like AIX on IBM machines, HP-UX on HP computers (earlier PA-RISC, nowadays Itanium based), Solaris on SUN/Oracle servers (SPARC), MacOS on Apple Macintosh machines). When the manufacturer designs and produces the machine and the operating systems as well, it can exploit the hidden features of the hardware in the operating system as well. This combination can bring a good performance. From version point of view the open-source operating systems are refreshed very frequently bringing mostly new features into the solutions. The vendor-based operating systems have support strategy for the version and patch management. Only bug fixes are provided between the planned version delivery dates

Almost the deepest level of the architecture is the persistent data store, the disks or storage. Here we have again the possibility to choose disks and build our own storage layer or buy a storage system from a vendor. Depending on the internal architecture of the storage among the others different behavior, feature set, and performance are provided. The storage attachment to the computing layer is a huge question because of the different aspects, like speed, reliability, flexibility, redundancy, etc. In version and component management we have to think about component exchanges (e.g. simple disk, disk type, or even controller), software updates or exchange as well.

The lowest remarkable layer is the network. If we think about network we calculate with active elements, like routers, switches, security control component (e.g. intruder detection systems, firewalls), access points, etc. The other part of the network is the cabling. The active elements are again combined software and hardware components; the wires (e.g. coaxial cable, optical fiber cable, twisted pair) are pure hardware elements with special plugs (endings). The version control of networks can be started at configuration changes, like modification of firewall settings, but the replacement of a cable or restructuring the network can also be mentioned as version problem.

III. VERSION CONTROL, QUALITY GATES

Everything is changing in the real world and in the IT systems as well. Our research goes into the direction of analog changes of the systems. It means that the required changes after sufficient test can be smoothly implemented into the corresponding layer without any stop, restart of the system or interference on the end-user side.

Currently the changes are collected and built into a so-called version, which can replace the previous state (or version). By implementing a version we are not really go forward analogously, but rather jump to the next version level in a discrete manner. Our goal is to design environments with more analogous versioning behavior.

The version control is almost in each case a system, which contains the documentation of the changes. It should be attached to adequate quality control as well. In our research we plan to extend this version control with recommended functions and procedures as well. The version control and version management is mixed many times. In our case we speak about version management, which is the mechanism in a software solution how it handles the versions to be implemented.

IV. SOFTWARE ELEMENTS

The applications, especially the business applications like ERP have more software layers. The end-user, who works on the user interface of the system, does not know anything about the system architecture, used hardware elements and network. But he or she should be always on the safe side, so they should not be disturbed. The main behavior of a system in his or her mind is the stability, which is represented by the followings:

- Same outlook of the user interface (only some enhancements are excepted)
- Availability (when they want to use, it is working)
- Acceptable speed (sometimes can be “slower”, but any increment in response-time is regarded as natural)

Depending on the distance from the end-user to the changes the end-user feels different the modification. E.g. some minor changes on the user interface can cause significant problem, but the whole exchange of the storage system could remain unknown. (Check Fig. 4.) The deeper level we work in an IT environment, the more redundancy is required.

As the overview in chapter II described in an ERP system or other business application not only the application solution contains software elements, but the hardware components contain as well software parts. The below list gives an incomplete overview on the usage of software in hardware element:

- User interface
- Configuration
- Internal logic (optional)
- Firmware

This kind of software elements are HW-dependent components and they are not redundant, but live once in a hardware element. The redundancy and availability cannot be managed on this software level.

The application level software components can have their own load balancing and redundant architecture. Unfortunately the version management is not able to work with this kind of features. The bigger systems have their own enhancement model. [3] These models offer different level of changes, enhancement of the software. Fig. 3 below illustrates the different level of software changes.

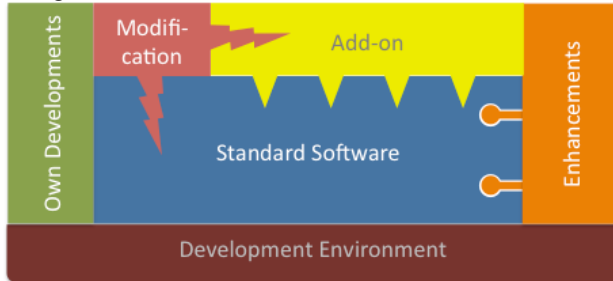


Figure 3. Software modification levels

During software modification the following four types can be significant:

- Configuration (not mentioned on Fig. 4, but can strongly influence the software behavior and outlook as well)
- Modification (changes of delivered standard components)
- Enhancement (soft changes of the standard software; these modification are implemented in predefined entry points without changing the original, only adding functions and logic)
- Add-on implementation (standard extension modules bringing new functionalities)

The own development is a special case, because it can be a totally separated program, applet, which does not harm the whole system, only gives new functions for some end-users.

Depending on the development strategy and model different approaches can be used in developments. One of these paradigms is the usage of so-called reuse elements. This reusability is a programming technique avoiding unnecessary thinking, designing, and programming efforts. If a task requires e.g. a special routine, it should be so designed and specified that other tools could it use later as well. (Well-defined signature, well-designed layers, etc.) Next time the routine should not be developed again to use it for other purposes. It brings less coding and faster development.

From version control point of view this kind of programming can be problematic if we should change some deeper object behavior, which is widely used in the software. Fig. 4 shows a similar dependency as described at the beginning of this chapter.

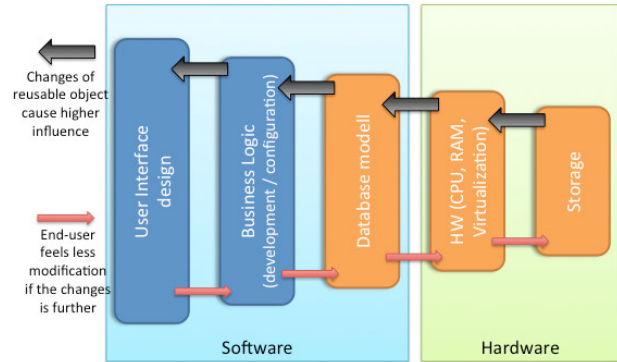


Figure 4 – Distance dependencies

The Figure 4 explains in one chart the dependency of influence

- On the end-user when changes are implemented on different level of the solution and
- On the software, when changes are implemented on different reusability level of the solution.

Similarities are conspicuous, because they effect in opposite direction.

Software can handle the version changes smoothly only if it has online version management. Under online version management I refer to a mechanism, which manages the running processes using current versions of programming units, objects, etc. in a manner, that the code (or object) versions are connected to the running sessions. After the new version is implemented the system records the new version, without removing the instances from the old ones. It keeps the running sessions using the old version alive. If a new session starts, it uses the already the new version. The online version management work strongly together with the session management, because it check when an old version is not used anymore and notifies the version management to deactivate the old version. With this online version management behavior the software components can be updated smoothly.

V. HARDWARE ELEMENTS, OPERATING SYSTEM, DATABASE MANAGEMENT AND VIRTUALIZATION

As described in the previous chapters an ERP system is not running alone, but some other components are needed as well:

- Database management system for ERP data store
- Operating system for running the ERP business logic machine and the DBMS
- Hardware elements, which run and manage the operating systems

Today's virtualization and cloud computing project a direction for us to follow. [6] In global thinking the virtualization can be followed, because we use it in many layer without checking the implementation (like network card teaming, RAID arrays for disks, etc.).

More important is the redundancy in the hardware layer, because it is the basis for virtual design. Our research goes in storage and database management systems redundancy and reliability as well. We expect that with the current available hardware elements we can build fast, reliable, redundant and cheap solution. For that own

defined storage nodes can be introduced with special file system offering. The Fig. 5 illustrates this idea:

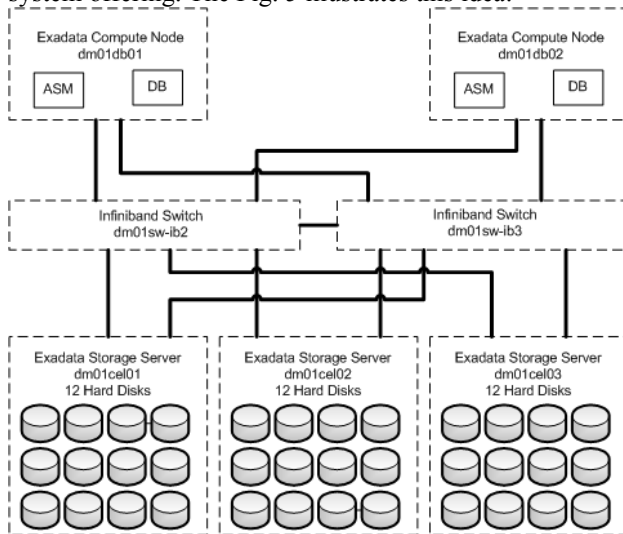


Figure 5 – Oracle Exadata architecture diagram

(Source: <http://blog.oracle-ninja.com/>)

I copied the Fig. 5 from the above-mentioned source, because it describes similar solutions with Oracle Exadata solution.

It is recommended to detach the applications from the operating systems, like in a high available cluster solution. It brings the availability and load distribution of the software higher level.

VI. CONCLUSION

This paper collects our current researches around building sustainable, stabile ERP environment using the current innovations.

I could figure out that the main point of the sustainability is the redundancy and the layer based exchangeability.

I have introduced and used the following definitions:

- Version control: not only a software managing the documentations of the changes, but rather

- Version management: software solution build in the application managing the software deployment and versioning
- Analogous versioning: continuous version update without stopping the system or components, and unknown by the end-user
- Online version management: mechanism to manage the old and new versions together in a software application

I have introduced a simple, not detailed mechanism, and procedure for adaptive versioning in this paper. These technologies and theorems can be used in education as well to design smaller application.

REFERENCES

- [1] A. Selmeci, T. Orosz, Gy. Györök: *Innovative ERP virtualization*, Intelligent Systems and Informatics (SISY 2013), IEEE 11th International Symposium, Subotica Serbia; ISBN 978-1-4799-0303-0 pp., 69-75
- [2] A. Selmeci, T. Orosz, Gy. Györök: *Potential of dynamic development in ERP environments*, 5th IEEE International Symposium on Logistics and Industrial Informatics (LINDI 2013) Wildau, Germany, ISBN 978-1-4799-1257-5
- [3] A. Selmeci, T. Orosz: *Modification free extension of standard software*, 12th IEEE International Symposium on Applied Machine Intelligence and Informatics (SAMI 2014), Herl'any, Slovakia ISBN 978-1-4799-3441-6
- [4] Orosz Gábor Tamás: *Integrált vállalatirányítási rendszerek. SAP: üzleti folyamatok és programozás*, ISBN 978-615-5018-52-7
- [5] A. Selmeci, T. Orosz: *Trends and followers in GUI development for business applications with implications at University Education*, 10th IEEE International Symposium on Applied Computational Intelligence and Informatics (SACI 2015), 21-23 May 2015, Timisoara, Romania, ISBN 978-1-4799-9910-1, pp 243 – 251
- [6] DANIELS, Jeff. *Server virtualization architecture and implementation*. Crossroads, 2009, 16.1: 8-12.
- [7] GOEL, Ms Shivani; KIRAN, Ravi; GARG, Deepak. *Impact of Cloud Computing on ERP implementations in Higher Education*. INSTITUTIONS, 2011, 5: 8.
- [8] HASSELBRING, Wilhelm. *Information system integration*. Communications of the ACM, 2000, 43.6: 32-38.

Experiences of Learning the Methodology of Development with Design Thinking in Multinational Online Course

Dóra Dancsics**, Péter Fehér*, Ádám Gazsi**, Dániel Marton**, Dávid Molnár**,
Bence Léki**, Levente Rádai**, Roland Varga**, Bálint Végh**

* Corvinus University of Budapest, Budapest, Hungary

** Óbuda University, Székesfehérvár, Hungary

dora.dancsics@gmail.com, pfeher@informatika.uni-corvinus.hu, gazsi.adam@gmail.com, martondnl@gmail.com,
dmoln@freemail.hu, leki.bence@gmail.com, radai.levente@amk.uni-obuda.hu, nfownl@gmail.com,
ballint999@gmail.com

Abstract — The paper presents the experiences of the participation in an multinational online course about the Design Thinking methodology. First it gives an overview the Design Thinking methodology in teamwork, and the details and results of the development project. In the second part, we analyze the way of thinking and use of design thinking originally and by our Hungarian group.

I. INTRODUCTION

The software development in team work can be supported with many project and team work methodology. One of them is the Design Thinking, which focuses on the start-up section of a project, the elaboration of the user needs. There are many types of Design Thinking methodology. Here we will discuss first the team work method, which we have tested in a 6 week long international pilot course of on openSAP platform for SAP open online courses. This course was available only for the invited institution, where we have performed a certain development project with a lot of other team for example from Belorussia, Germany, USA, United Kingdom, China, Nigeria, Turkey, Taiwan and from many other countries.

From week to week we passed through the phases of the methodology and presented the results for the learning community and at last we rated our results among three other groups. During these phases we have sensed, that our group is thinking a little bit on different way, so after we presented the Design Thinking methodology and the pilot project, we try to look after the reasons for this difference by analyzing the team results.

II. THE METHODOLOGY OF DESIGN THINKING IN TEAMWORK

A. Design thinking process

In the Design Thinking process there are five main steps, which starts with **Empathy** phase to gain insights on the needs of the audience of the development subject, and goes forward to the **Define** phase, where we have to use the insights of Emphasizing results to identify a problem to focus on. Next step is **Ideate** phase, where you take that problem, and solve that with many unexpected ideas with using ideation-techniques.

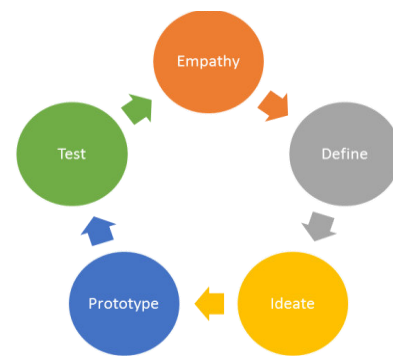


Figure 1. The Phases of development in Design Thinking Team Work Method

To generate options to solve the problem what we can come up with we create several prototypes in the **Prototyping** phase. It is important, that we should create a raw experience for us and for our “audience” to learn here, what works and what does not, and these prototypes should be tested in the last phase **Test**. We needed to learn a different way of thinking in teamwork. And it is also an important

Generally, when we start a project, all of these key values are empty, for example, we did not have a lot of empathy for people using the subject of the development, we did not really know, what the problem is, and we did not have a lot of ideas, we have not so much test result yet. The goal of design thinking is that it can give a framework to solve such hard development tasks, where the previously listed features are present, and after we work with them day by day, we got a few insights about the needs. We had some ideas and test results of some prototypes.

According to this, during the course each time we work on our project these key points become clearer and after 6 weeks, we can say that, we successfully developed a new and innovative solution for the user needs. Moreover we needed to learn how to make decisions as a team, and how to work together in virtual working space, because we could work mostly at night.

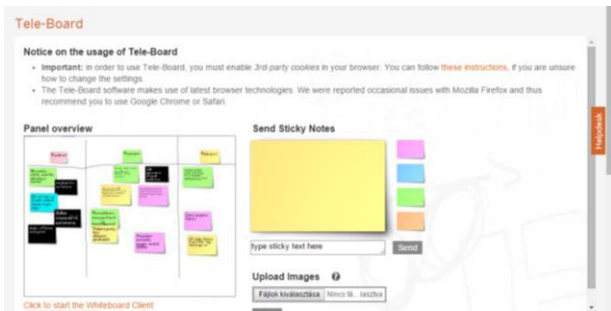


Figure 2. Tele-Board for virtual meetings (SAP® tool)

So we learned how to use this process for our own way of action in virtual meeting place using video conference tool and a Tele-Board platform, which is a shared virtual white board.

B. Iteration process

This whole process - with the mentioned 5 steps - works, if the process is iterative. The iterative cycle helps to learn deeper insights the fastest, get more and better idea and based on that, refine the prototype. We have to try to go through the cycle as much as we can in the available time. For example, in a professional project, if we have 10 weeks, we had to try to run our cycle 10 times, but if we have only a few weeks, had to try to run our cycle also 10 times.

In iterative cycle, we have 2 moving targets: need and solution. When we come around the block and come up with the definition of the need, come up with some ideas, we test the prototypes, we had to learn if we identify a wrong need, and we need to change the result of the Define Phase.

Next, we can have the right need, but with a wrong solution, therefore we need to change the solution. This is the hardest part of the design process, but helps to identify the right need of the user and the right solution. However both of the need and the solution are moving targets, through the iterations, they come closer to each other.

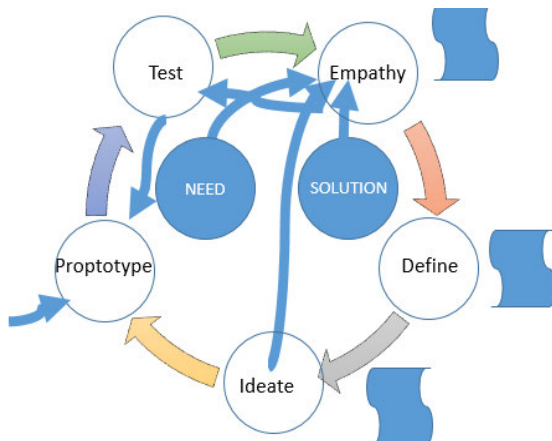


Figure 3. Iteration ways in Design Thinking

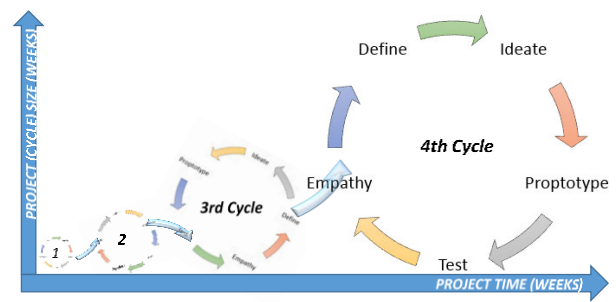


Figure 4. Mileage in iteration [1]

Generally, a Design Thinking process cycle is starting with Empathy Phase to come up with the user need, but we can change the steps if we have insight and idea to build a prototype and test. But we need to be careful because it has a risk to choose the wrong direction, so we have to go around the block as many times as possible to reach a refined, right user need and solution based on a deep insight. Furthermore, in the iteration process, we can use the technique of Mileage. With the iteration technique, we can learn the fastest, if we have more and more cycles. But we can start with a short cycle within the iteration process and we can give ourselves a chance to have as many projects as possible, but each after each other can be larger and larger cycles.

III. THE PROJECT OF DEVELOPMENT OF SERVICES AND FUNCTIONALITY OF VENDING MACHINES

Our project was to redesign the working of vending machines. We had to start with the Emphasizing. As we previously mentioned, we do not need to start with the Empathy, but it is recommended to understand our end users and their needs.

A. Emphasizing

In this phase we had to gain insights on the needs of the vendor machine users. For that, all of us had to visit individually a vending machines or coffee machines or drink machines, and asked some users, who visited the machine right in that few teen minutes about four aspects of use of the machine:



Figure 5. The result of the Empathy phase

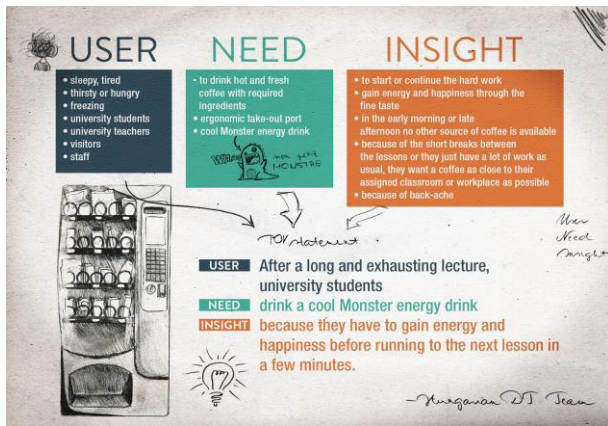


Figure 6. The Point-of-View selected to provide solution for it

- What is their actual life situation?
- How and why are using this machine?
- What do they struggle with?
- What do they wish from the machine?

After the survey, we had to summarize the group findings. The merged result is shown on Figure 5.

B. Define the Point of View to serve with solutions

In the second phase we had to define only one chosen Point of View (PoV in short) of the vending machines users based on the key insight about them. It is important that the PoV has to be specific not general. First we analyzed the collected data of insight, and a created individual User+Need+Insight statement.

Then we tried to merge and combined the individual statements by looking after similarities and differences and correctness, as it is presented in Figure 6.

C. Ideation

The third phase of design thinking is the Ideation phase. In this phase we had to generate a lot of ideas and concepts individually for our vending machine development through Brainstorming technique. During the idea collection, we had to focus on the PoV. In the second step, together as a group, we had to sort the ideas into most radical, the safest and most delightful idea groups.

For this task we applied the Tele-Board, which result is presented in Figure 7 of.

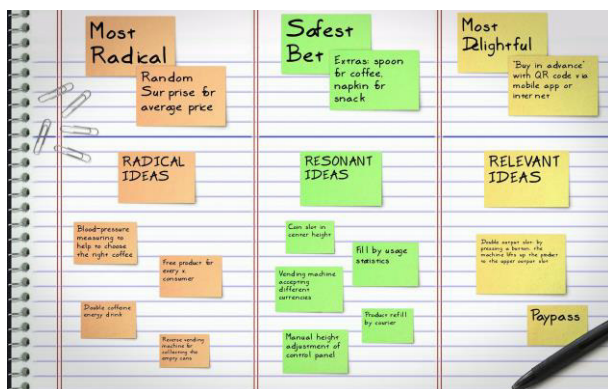


Figure 7. Results of Ideation Phase in three categories

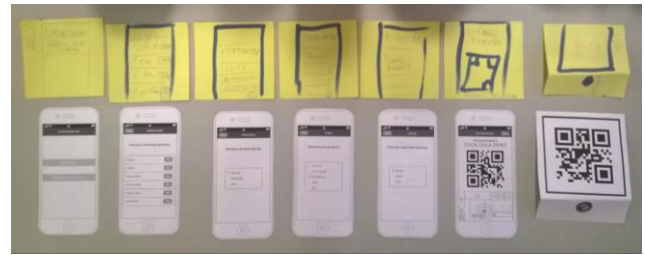


Figure 8. The rough and the refined prototype

D. Prototyping

Next week we reached the fourth phase and we had to make some prototypes for our selected ideas. We had to make “ten-minute prototypes” to test with our end users. They could give us a relevant and honest opinion of it and even if they not satisfied with our prototype we just can throw it away and search after other prototype or idea.

E. Testing

The last phase of our progress was the Testing, which was really delightful. We could see how people react to our work and for our new ideas.



Figure 9. The testing phase with real users

IV. DESIGN THINKING IN INDIVIDUAL WORK

In comparison with the Design Thinking in team work, we can use the Design Thinking the individual development projects, such as designing user centric components of software system, like user interfaces. During the development, the keys are very similar, but the Empathy and Define phases has got its own iteration cycle. (Figure 10)

THE DESIGN THINKING STEPS.

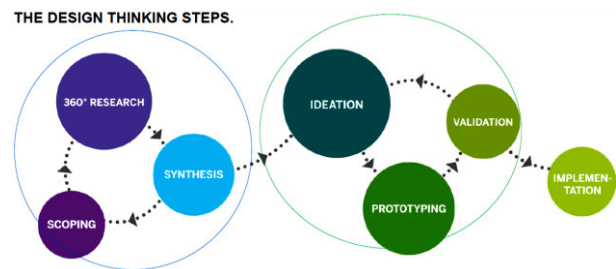


Figure 10. The Design Thinking steps with double iteration cycles (SAP® tool)

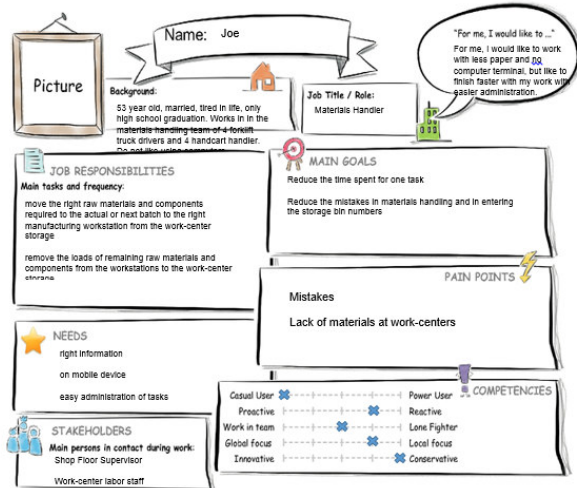


Figure 11. Persona (SAP® tool)

To solve the task of 360 degree research first, we have to create the story of the user to get insight to the user's life, work, role and skills. Second step is to create the persona of the user based on the collected information to see his/her goals, skills, mistake opportunities, needs, responsibilities and life situation. (Figure 11)

Third step is to create the user experience journey to collect the user actions (steps) and his/her mindset, touchpoints and pain points during the actions (Figure 12.), which can help the scoping phase. If the result of Scoping right, we can proceed to the next process cycle.

V. DIFFERENCES IN THINKING IN DESIGN PROCESS

The significant feature of Design Thinking – focus on the real customer needs and deep insights and recheck it again and again – is a different way of product development. Normally, we start with the analysis of the problem or the needs, but the way of design thinking emphasizing suggest going to the field of use of the project subject by the users. This can help in any phase to counterbalance our instinct to jump right to a solution.

During the phases from week to week, our team had tried to follow the instructors' recommendation, but we could

sense that in every phase there was the effect of the anticipation or intuition from our mind: what will be the result, the solution and how can it be connected to the results of the present phase. Sometimes we forgot not to think forward as in a chess game, and this could break the focus on the actual task. For this opportunity, our course mentors were helped the work of the teams, and they had recommended to rethink some points of our results, especially the Point-of-View and the Ideation map.

After the course closing we discussed in the team, that we just envisioned our ideas during the Empathy phase, although these idea were not so clear than after the Ideate phase. Therefore our focus on the Point-of-View was not so strong and we lost it a little bit during the Define and Ideate phase. But our mentors helped to refine it, so some iteration steps were included into the design process.

In the Prototype phase this focus was regained, and our prototype was rated almost to 100%, which shows, that our thinking was/is very solution-oriented, and therefore this methodology can really help not only to develop a good product, but answer the right need f the user with the right solution.

VI. SUMMARY

In this paper, we presented briefly the methodology of Design Thinking in team work learned and practiced in a pilot course of openSAP platform and discussed the steps of the performed development project. We analyzed the experiences during this development project and discussed the benefits and strong points of the methodology. And for an outlook we compared the Design Thinking in team work and in individual work.

ACKNOWLEDGMENT

We would like to say special thanks to openSAP for inviting us to this very exciting and innovative course.

REFERENCES

- [1] openSAP Design Thinking Pilot Course, URL: <https://open.sap.com/courses/dt1-pilot2>, 2015
- [2] openSAP Build Your Own SAP Fiori App in the Cloud, URL: <https://open.sap.com/courses/fiux1>, 2015

Duration of the Journey: less than 1 min + duration of materials handling

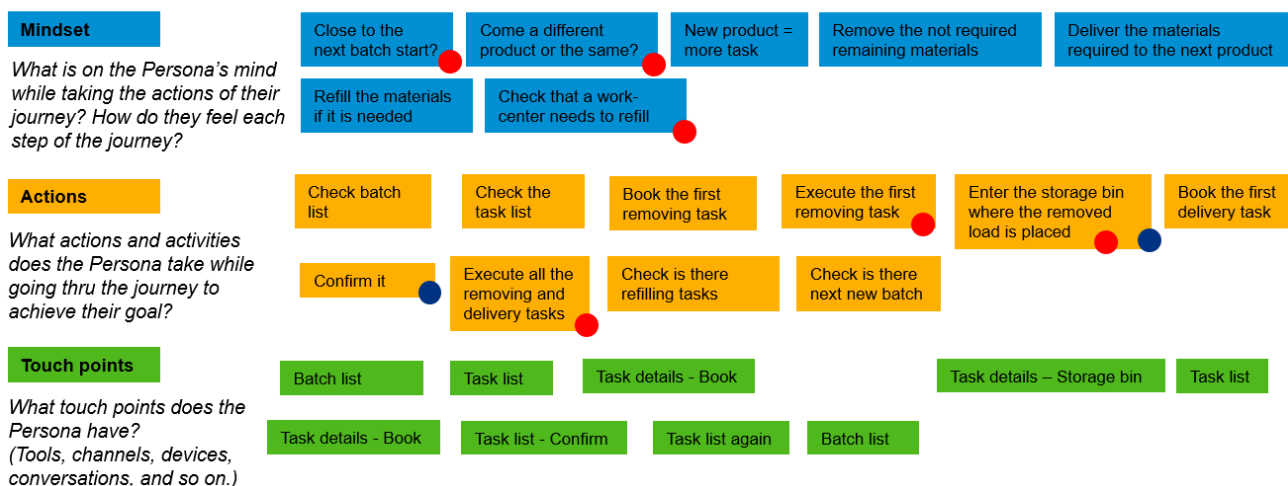


Figure 12. User Experience Journey of a plant material handling scenario (created by SAP® tool)

Examination of Eco-environmental Relations of Urban Areas

Péter Udvardy*

* Óbuda University/Alba Regia Technical Faculty, Székesfehérvár, Hungary
udvardy.peter@amk.uni-obuda.hu

Abstract—The main reason for the examination was to clarify the interactions between urban areas and its surroundings. To reveal the conflicts in urban areas measurements were made around urban and semi-urban areas considering air, water and soil status. The protection of the urban green areas and the prevention of the built structures had a top priority in this research considering the increasement of the urban environmental life quality and the development of the environmental protection infrastructure. Therefore it was necessary to study and evaluate the so called urban ecological parameters. The next step was the creation of an environmental cadaster and the elaboration of an eco-friendly urban development plan proposal. The results of these complex studies should be a part of the urban development processes in the near future.

I. INTRODUCTION

Landscape ecology which includes urban ecology is a multidisciplinary subject since it exploits several results of social and natural sciences. The location of urban areas was basically influenced by environmental components. The relationship is mutual because not only the environment defines the urban areas' location but the urban areas have an effect on the environment. There is a permanent flow of material, energy and information between the cities and the environment. Although the two components had been examined separately several times there is only a little knowledge about their impact on each other especially the ecological aspect. The aim of our study is a better understanding of these relationships.

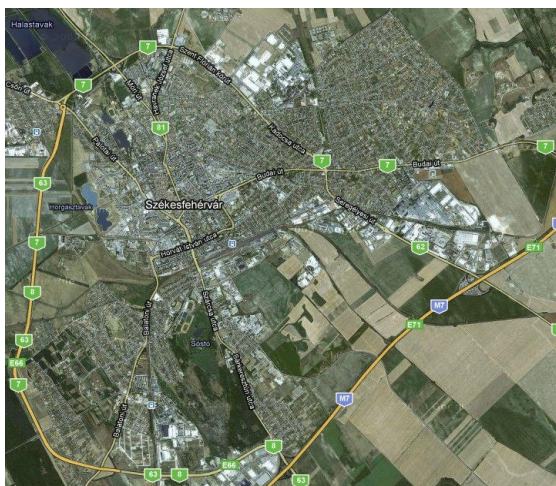


Figure 1. Map of Székesfehérvár

source:maps.google.com

Székesfehérvár lakónépességének száma

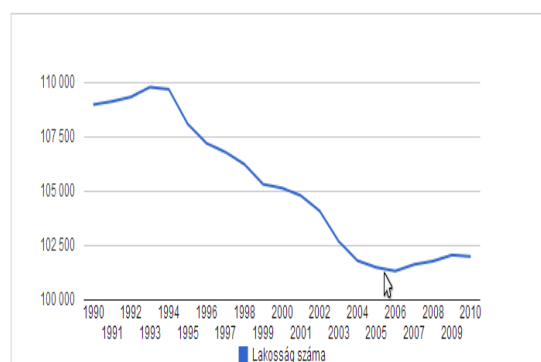


Figure 2. Population changes in Székesfehérvár

During our work the connection between the urban and periurban areas and their mutual effects on each other were studied. Important part of our work was the examination and evaluation of the so called urban ecological parameters (climate, water balance, vegetation, noise, air pollution, etc.); creation of an environmental cadaster and working out of ecological based suggestions for urban management and city development. These ecological city development suggestions become an organic part of urban development and urban management programmes.

The examination of the urban environmental conflicts and their location aroused as a major task of the project or rather the decisions connected to the urban land use. The spatial elements concerned may consist natural systems such as agricultural and natural ecosystems, land coverage elements, ecological networks, and green areas. The maintenance of the green areas and the avoidance of damages caused by erosion and drying in built environment is an important issue which can be achieved by the improvement of environmental quality, environmental protection infrastructure, and protection of built and cultural heritages.

Figure 1 and 2 show one of the studied cities' map and its changes of the population.

II. DATA AND METHODOLOGY

The main objective during the land ecology examination data collection procedure was the complexity of data sources which assure the high level data procession. During our work satellite images, topographic maps, meteorological datasets, green areas maps, etc. were collected and used. Some of the data were available in raster format but most of them were used and examined in vector format such as point located spatial data.

Long term statistical datasets which were collected by the Hungarian Statistical Office (KSH) (i.e. population, precipitation, temperature, wind, etc.) were also used for timeline analyses. Detailed study was made in Székesfehérvár, but only major urban ecological data were examined in Szombathely.

The study of changes in demographical data of a city considering the urban environmental examinations is very important since the population carrying capacity of a settlement can be seen. Demographical trends correlates strictly with the intensiveness of environmental usage and environmental impact (Konkolyné, 2003).

The population of Székesfehérvár decreased permanently in the past decades, the number of inhabitants is around 100.000 persons. This number correlates with the general trends in Hungary, some of the inhabitants move to the periurban areas. The population density is around 600 persons per square kilometer which is not too high compared to the western European countries. The advantages (i.e. green areas) and disadvantages (less developed areas) of the city must be revealed by the city's decision makers in order to make the city more comfortable. The creation of spatial usage map and spatial structure map is a good tool in decision making process.

Different amount of samples were collected according to the environmental elements and affecting factors in the examined cities. At the end of the project the following matrix is used to define the given status of the environment. Table 1 shows the environmental matrix.

TABLE I.
ENVIRONMENTAL STUDY MATRIX

Pollution source	Concerned environmental elements, and effect factors					
	Air	Water	Soil	Green surface	Waste	Noise
Traffic	X	X	X	X	X	X
Industry, trade	X	X	X	X	X	X
Health	X	X		X	X	X
Wastewater management	X	X	X		X	
Agriculture	X		X	X	X	
Cities and built up environment	X	X	X	X	X	X

During our examination the advantages of GIS software were utilized. The IDRISI Taiga software showed the directions of the urban areas' increasement. The ArcGIS 9.3 software was used to visualize the spatial structures of

the urban areas. The two software above complete each other during the data processing procedure and can handle

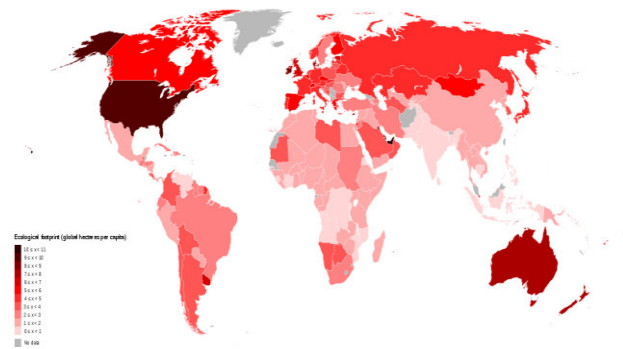


Figure 3. Ecological footprint of the world

source:Wikipedia.org

both type of data models (raster and vector data).

The modelling of the urban areas' increasement was done by visual interpretation, the spatial structure analyses was made by the help of GIS software.

III. RESULTS

The spatial structural examinations were conducted considering several aspects. One of these aspects was the

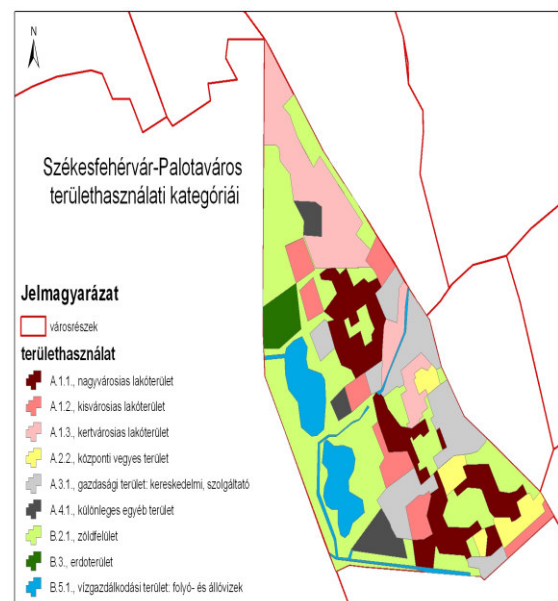


Figure 4. Land use categories in Székesfehérvár

demarcation of the green areas which was a great job since there are big green areas in Székesfehérvár. The green areas which are intended to increase the ecological capability of a city must be integrated to the whole city's spatial structure (Konkolyné, 2003). The large part of the urban dendroflora (trees) is not native at that area. The number and species of the urban trees are various considering the different cities and a large number of different tree species can be found in the cities (Nagy 2008).

In this study the protected trees of the different cities were localized during the examination process of the green areas. As it is written in the literature it can be said that significant number of trees can be found in alleys and city parks. Protected areas and parks were localized in a different category. Figure 5 and 6 show the location of trees and green areas of Székesfehérvár.

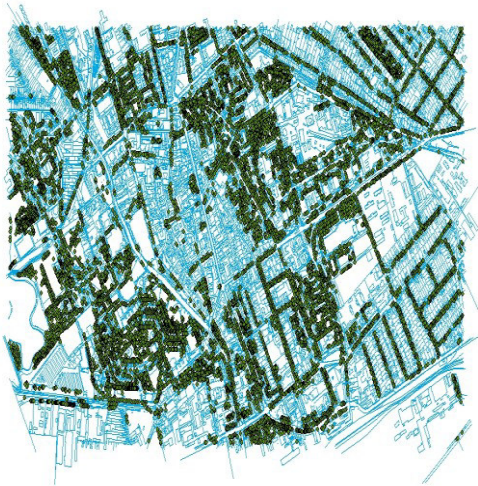


Figure 5. Tree map of Székesfehérvár downtown

Land usage index is a complex indicator which characterizes the urban structure. Thus land use types were created which were characterized in details by further indicators which consider their ecological contents. With the help of the analyses of the existing databases available at the local municipalities the land use types can be interpreted easily thus the studying of the urban ecological structure also can be made. On the basis of the existing database a new urban area use database can be created which considers the ecological aspects and helps to make a complex comparative study with other cities (Nagy 2008).

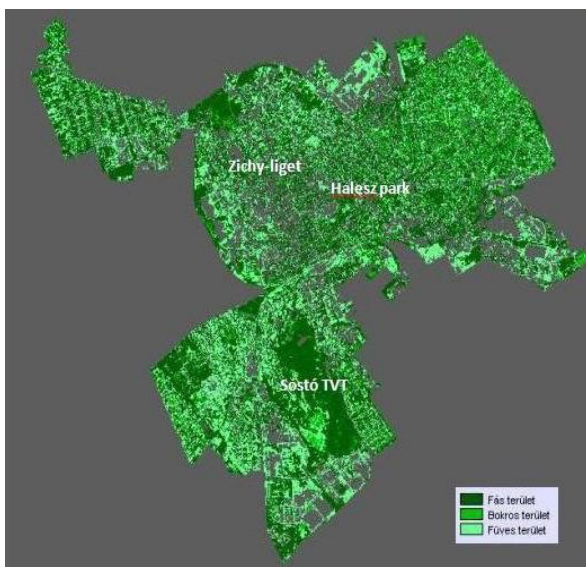


Figure 6. Green areas in Székesfehérvár done by using remote sensing techniques

Land use changes and urban area increase were examined in Szombathely and statistical analyses were done. It can be seen that the city area doubled in the last 200 years which was caused by its favorable location, geographical conditions and special natural-economical potential.

The increase of the urban area in Szombathely is not parallel with the changes of the population but because of the increase of the industrial areas and the need for residential areas. The size of the residential areas get bigger in the past decades as it happened in west Europe (Mizseiné et al., 2012).

One of the main objectives during the study was to determine the urban ecological footprint. The measurement method for the environmental impact was developed by Wackernagel and Rees called 'ecological footprint' evaluation method. This value is determined in hectares and contains the input and output energy and materials used by the certain group of people. Figure 3 shows the ecological footprint of the world.

After finishing the calculation the result reveals the quantity of land and water surface that is necessary to sustain the system. The impact of any regional economy or enterprises or sport event or a single human being can be specified with this method. Figure 4 and 7 shows the land use and spatial structure of Székesfehérvár.

The results of studying the human-environmental interaction in urban areas help the determination of real and specific environmental problems and the specification of the environmental sensitivity of the groups. In former researches on individual and its surrounding environment scientists mainly focused on social environment instead of natural and material environment. These factors were only studied in the recently.

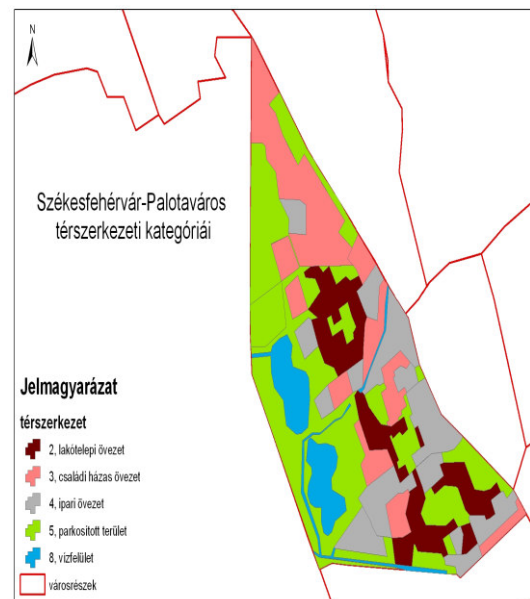


Figure 7. Spatial structure in Székesfehérvár

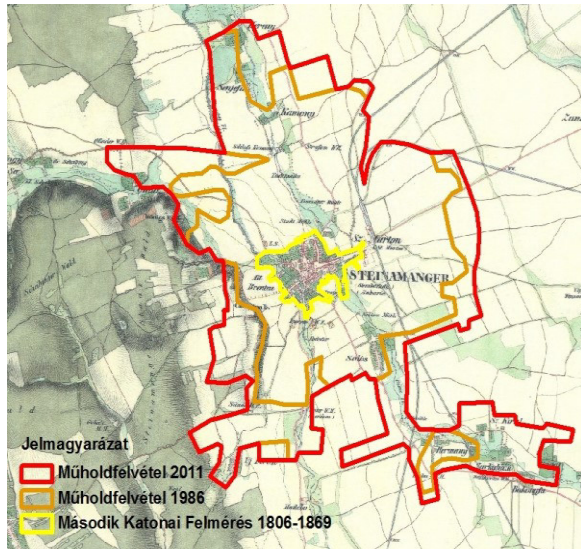


Figure 8. Land area changes in Szombathely in the past 200 years

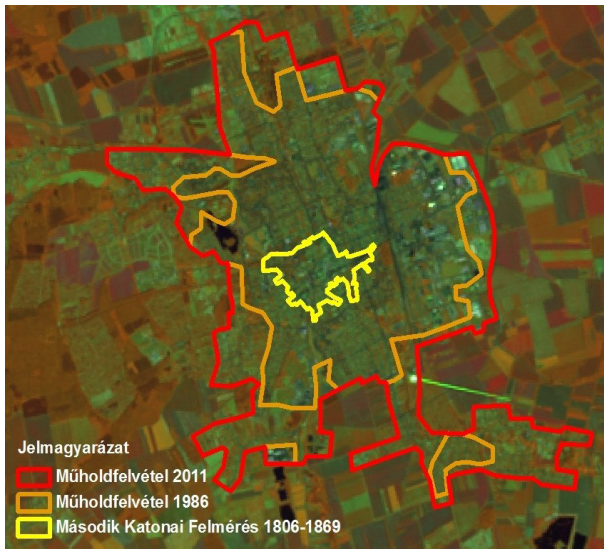


Figure 9. Land area changes on LANDSAT image

The urban ecology science reveals the connection and cause and effect interaction between the urban inhabitants and the biosphere, the environmental harms and conflicts and highlights the regularities of social mechanism and social and psychical reaction (NAGY 2008).

Figure 8 and 9 shows the land area changes of Szombathely in the last centuries.

IV. SUMMARY

The aim and objectives of the research study is in harmony with the cities development projects thus the results can be used in decision making process. The implemented project proposals can be specified upon the results of the health check and the professional and NGO's' requirements. This study gives a scientific support and adequate information for decision makers on these specific fields: urban environmental hygiene, precipitation sludge management, sewage sludge management and purification, communal waste management, local public transportation management, air pollution management, drinking water supply, energy management, green area management, wreck control, natural and built heritage protection, urban environment protection, landscape protection, biosphere and natural conservation. management, local public transportation management, air pollution management, drinking water supply, energy management, green area management, wreck control, natural and built heritage protection, urban environment protection, landscape protection, biosphere and natural conservation.

REFERENCES

- [1] KONKOLYNÉ, GY.É. (2003): Környezettervezés. Mezőgazda Kiadó, Budapest, 398 p
- [2] NAGY, I. (2008): Városökológia. Dialóg Campus Kiadó, Budapest-Pécs, 335 p.
- [3] MIZSEINÉ, Nyíri Judit – HOROSZ-GULYÁS, Margit–UDVARDY, Péter –KATONÁNÉ, Gombás Katalin– KATONA, János: Complex eco-environmental study on urban area of Székesfehérvár. Városok öko-környezetének komplex vizsgálata, szerk.: Albert Levente, Sopron2012. ISBN978-963-334-084-4

BIG GEOSPATIAL DATA PROCESSING IN THE IQMULUS CLOUD

*Angéla Olasz^a, Binh Nguyen Thai^{b,a},
Roberto Giachetta^{b,a}, Dániel Kristóf^a*

^aInstitute of Geodesy, Cartography and Remote Sensing (FÖMI), Budapest, Hungary

^bEötvös Loránd University (ELTE), Faculty of Informatics, Budapest, Hungary

ABSTRACT

Remote sensing instruments are continuously evolving in terms of spatial, spectral and temporal resolutions and hence provide exponentially increasing amounts of raw data. These volumes increase significantly faster than computing speeds. All these techniques record lots of data, yet in different data models and representations; therefore, resulting datasets require harmonization and integration prior to deriving meaningful information from them. All in all, huge datasets are available but raw data is almost of no value if not processed, semantically enriched and quality checked. The derived information need to be transferred and published to all level of possible users (from decision makers to citizens). Up to now, there are only limited automatic procedures for this; thus, a wealth of information is latent in many datasets. This paper presents the first achievements of the IQmulus EU FP7 research and development project with respect to processing and analysis of big geospatial data in the context of flood and waterlogging detection.

Index Terms — Distributed computing, flood detection, Big Data, raster data partitioning,

1. INTRODUCTION

The four-year IQmulus project (<http://www.iqmulus.eu>), funded by the 7th Framework Programme of the European Union, is targeting to enable optimized use of large, heterogeneous geo-spatial data sets (“Big Geospatial Data”) for better decision making through a high-volume fusion and analysis information management platform. An end-to-end involvement of users is ensured through the implementation of two concrete “test beds” (Maritime Spatial Planning & Land Applications for Rapid Response and Territorial Management) to show the benefits of the approach. The contribution of large number of users from different geospatial segments, application areas, institutions and countries have already been achieved. The consortium consist of partners representing numerous different facets of the geospatial world to ensure a value-creating process requiring collaboration among academia, authorities, national research institutes and industry.

User requirement collection along with scientific and technical state of the art analysis was carried out in the first phase of the project to identify relevant development directions, in which IQmulus can yield significant improvements. The project starts providing concrete results in terms of processing services and system architecture. Services consist of algorithms and workflows focusing on the following aspects: spatio-temporal data fusion, feature extraction, classification and correlation, multivariate surface generation, change detection and dynamics. System integration and testing is currently being done. The IQmulus system will be deployed as a federation of these modular services, fulfilling the needs of the above-mentioned scenarios but also suitable for the construction of further solutions thanks to the modular approach.

2. METHODS AND MATERIAL

According to IQmulus specification, services should run on distributed environment. The Apache Hadoop (<http://hadoop.apache.org/>) open source framework has been selected as a basis for architecture development. This framework offers the execution of large data sets across clusters of computers using simple programming model known as MapReduce. It is supplying the next generation cloud-based thinking and is designed to scale up from single servers to hundreds of workstations offering local computation and storage capacity. Thanks to this solution, users can achieve rapid response to their needs used by their own datasets in cloud processing [1]. However to take advantage of today’s state of the art computing, previous data processing methodologies and workflows have to be revisited and redesigned. Actual deployment is based on a distributed architecture, with platform-level services and also basic processing services invoked from clusters and conventional computing environments.

2.1 Distributed File System (Hadoop Eco-system)

Data have become a part of everyday’s work. Depending on application areas, some of them need high data availability; others need high computational functionality on large data sets. Scalability of any framework is a crucial point of

success or failure, however not everyone can afford large hardware investments in a short period of time, not to mention the processing power is not increased proportionally with hardware investment.

Hadoop is an open source framework/eco-system designed for capturing, organizing, storing and sharing, analyzing and processing data in a clustered environment giving the solution for linear scalability of both storage and computational power. The ability of storing and analyzing large amount of data on clustered computers reduces the invested cost remarkably [2].

Primary approach of Hadoop: “Data should be distributed and minimizing network latency by moving algorithms across clusters in a distributed manner”. Hadoop have the following properties:

1. Data sits on one or more machines.
2. All nodes contain both distributed file system and data processing capability.

Hadoop is completely modular. Its core components are HDFS and MapReduce which is required; others are optional depending on your need (Figure 1). Some of the widely used components are HBase, Hive, Pig, and SpatialHadoop. Hadoop distributed file system have been designed to store text-based data. Incoming data are split into smaller chunks, namely data blocks which are 64MB by default, this limit have been increased to 256MB from version 2.x [3].

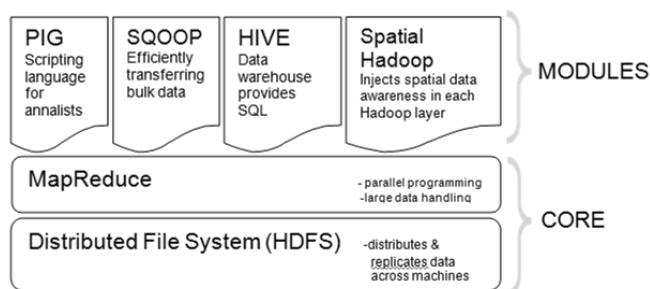


Figure 1. Architecture of Hadoop

Originally, HDFS is designed to store text-based data, however a large portion of our geospatial datasets are stored as structured binary files. Hence, data partitioning over HDFS had be solved in the first place. On the other hand, services are implemented in different languages varying from script to object-oriented languages, for example LIDAR processing services are implemented in C++, general processing services are implemented in Matlab, Java, Visual C++ and remote sensing algorithms related to land-based scenarios are implemented in C#. Therefore, the distributed framework has to support a wide range of runtime environments.

On architecture layer, we have to find solutions for data partitioning and supporting a runtime environment for services written in different programming languages.

2.2 Data Partitioning

To process raster data in distributed environment data decomposition is crucial. In distributed computing environment access data in a certain location, decompose to smaller chunks ('sub-raster blocks') to achieve optimal fast processing, collect the result of the algorithms from processing units and rebuild resulting data. Many approaches exist with different pattern of coordination and communication between the processing steps and results. The main aspect is how to decompose input raster to smaller partitions. A number of alternative methods exist to raster data partitioning as well. Single pixel algorithms are processed for each pixel in isolation. Defined extent algorithms are processed for each pixel in a context of a pre-determined number of surrounding pixels. Region-based algorithms are considered to apply in a geographical application where a greater neighbouring area has significance where the study area is grouped into homogeneous segment or stratum. The processing algorithms and parameters are designed for every homogeneous strata [4]. This stratum considered as a landscape unit which is homogeneous in aspect of physical geographical factors, or any targeted application area.

One of the most trivial ways to decompose data is to spatially split the original data into smaller data chunks, which can be processed independently on processing units. However decomposition methods are tightly related to processing methods and the number of processing units available in the distributed environment.

As a proof of concept, we have developed a GeoTIFF decomposer application using GDAL library. Decomposer takes a source file or directory, looking for GeoTIFF files and decompose them into smaller partitions. It takes two parameters, namely a predefined grid size (N x N) and number of processing units available in the distributed computing system. The number of processing unit is needed later for data distribution. After decomposition has been successfully performed, decomposed data are being uploaded to computation nodes. Within the next year, we will work on more sophisticated splitting algorithms to suffice user stories and showcases gathered in user requirement analysis phase [5].

3. RESULTS

3.1 Development of processing services

IQmulus develops functional and domain processing services in order to: maximize the use of big geospatial data, provide task-specific packing and delivery of data sets, support data quality evaluation, and provide support for

analyzing quickly changing environmental conditions. Services are developed to fulfill requirements related to data preparation, feature detection, change detection and presentation/visualization of big geospatial data.

The specific services and solutions put in the focus of the current presentation have been developed on the basis of the AEGIS open-source framework, a geospatial toolkit developed in a joint work between the Institute of Geodesy, Cartography and Remote Sensing (FÖMI) and the Eötvös Loránd University (ELTE), Faculty of Informatics [7].

3.2 Data management in the cloud

To enable geospatial processing in the cloud, the spatial data is preferred to be stored in the distributed file system. The storage forms are general spatial file formats (such as Shapefile, LAS, or GeoTIFF) to enable the direct consumption by most geospatial toolkits [8]. It must also be taken into account that large files are divided by the file system into smaller blocks before distribution to enable the parts to be processed individually. These blocks should be separately readable by the libraries, requiring all blocks to have a readable format. In contrary, the general partitioning methodology of HDFS does not take the file structure into account when splitting the input file.

A custom partitioning solution is provided, that performs partitioning and allocation of the datasets in HDFS using predefined strategies. Instead of the file blocks originally created by HDFS, these constructed items are individual files that can be interpreted element-wise. Strategies can be based on spatial extent, spectral space (in case of imagery) or any other property. Some strategies may have special purpose, e.g. creation of a pyramid image for visualization.

3.3 Processing framework

For implementation of the flood detection services of the IQmulus project (considered as Land Showcase 3) the AEGIS framework was used. The AEGIS framework is geospatial toolkit for both vector and raster data [8]. It is based on well-known standards and state of the art programming methodologies [9]. Due to its extensible infrastructure, it can support a variety of execution environments, including distributed platforms. A built in execution environment is responsible for handling algorithms, and managing the execution over the designated platform. In case the execution platform is Hadoop, AEGIS transforms the operations to use the MapReduce scheme. The MapReduce paradigm can be considered as the Map phase performing the initial operation, whilst the Reduce phase is responsible for post-processing the results [10, 11]. Thus, operations can be adapted using the following process:

1. When executing an operation based on specific spatial extent or other properties, the required files are selected using indices and the metadata catalogue.
2. Based on the specified method and input data, the operation is determined from the operation catalogue.
3. The operation and input are forwarded to an operations engine working over the Hadoop environment. Each input geometry is processed in parallel by separate MapReduce tasks.
4. If required, the results of the operations applied to multiple geometries can be merged using post-processing.
5. The result is written to the distributed file system by the Hadoop environment.

Post-processing is a key step, which is usually a kind of merge operation performed on the initial results. In many cases there is no need for post-processing. For example, binary thresholding of an image can be performed element wise, thus when processing partial images the results are independently valid without the need of any merging.

Some operations may require a special approach and cannot be applied directly. One such case is histogram equalization, which can be performed in two steps. The first step computes the histogram of the individual image parts in the Map phase, then merges the histograms and computes the mapping of the values in the Reduce phase. In the second step the mapping is applied to the parts using the Map stage, resulting in the transformed images. Thus, no Reduce function is required for the second part. Due to these circumstances, it cannot be stated that all existing algorithms can be applied directly in the MapReduce environment without any prior investigation, but the required additions and development are marginal compared to complete reimplementations.

Nevertheless, there is a need for an operation environment enabling the execution of algorithms as Map functions and performing optional post-processing as well, which is an operations engine run by AEGIS [7].

2.3. Workflow definition and execution

Domain-specific Languages (DSLs)

Concerning workflow definition and execution, the concept is to develop a set of domain-specific languages (DSLs) that make the definition of several types of spatial data processing and analysis tasks simpler. DSL can be very close to natural language, which is descriptive and expressive. A well designed domain specific language can be very expressive, useful and understandable even for non-software developers. However, modeling a natural language or processing natural languages is not an easy task; therefore we have to look for existing programming languages with very high descriptive and expressive power [6]. Workflows defined in a DSL can also be compiled to run on parallel processing environments. The compiled algorithms can be

packaged as declarative services to run on the respective environment they have been compiled for.

The Workflow Editor

The Workflow Editor is a web-based application using domain specific language to define geo-processing services and working dataset developed in the framework of the IQmulus project. The actual look of the workflow editor can be seen on the Figure 3.

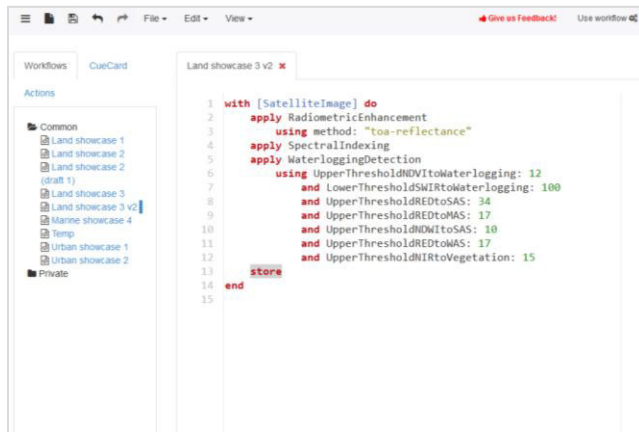


Figure 3: IQmulus workflow editor

On the figure the high-level DSL can be read for the Land showcase 3 which is a complex image processing workflow for waterlogging detection. The workflow implementation was originally done by GIS-experts and formulated using the DSL (which can be considered as in other solutions graphical modelers of a geo-processing workflow creation). Language designers have been analyzing the user requirements documentation, collecting all the keywords (nouns, verbs and adjectives), data types, service names, actors in order to create a language that meets user requirements in the respective domains.

Workflow designer is providing two user modes, for GIS-expert users and end-users. DSL user interface for GIS-experts provides extended language elements giving more control for experts over data sets and services, whereas for end-users a limited, more compact and lightweight version of DSL have been provided [7]. In this approach there is possibility to avoid very technical parameterization considering the user's needs, although an expert user has right also to change the previously implemented metrics. In the future vision on the CueCards section all the required information are provided to support the usage of the Workflow Editor. After the workflow is selected and edited by the user and also the dataset is selected for the appropriate workflow the process need to be executed. The process goes through on three main steps until it reaches the services and the data in the cloud. The parser, the interpreter and the job manager are the three main components.

2.4 Visualization

Visualization techniques are always playing important role to communicate the results to the end-users of a project. Even in the geospatial world we can consider the applied visualization support have more efficient impact in the dissemination of the results. However, in the traditional GIS systems the next generation graphical solutions are not implemented. IQmulus project stated the goal to “leverage the power of modern graphics processing units (GPU) technology” in an application of the GIS domain. The targeted user groups of the project defined in two levels: expert users and decision makers/non-experts. In the implementation it has been also considered, Fat Client (IFC) visualization software as the playground of the experts and the Thin Client as the interactive visualization support for non expert users. The architecture of the IQmulus visualization can be seen on the figure 4. On the figure the commutation between the implemented modules of the project are also visible.

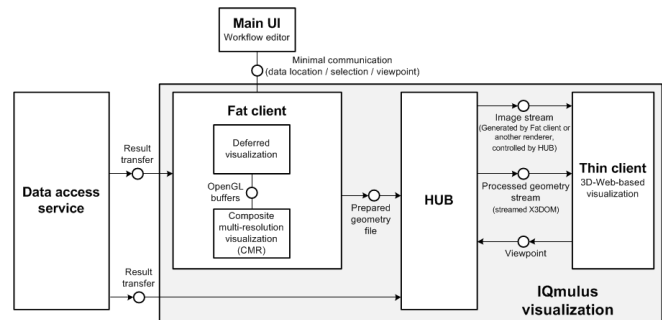


Figure 4: Architecture of the IQmulus visualization [12]

The Thin client is a web based 3D viewer for different geospatial data which are stored in the IQmulus cloud. The accesses of the results are realized by the automated visualization adaptation framework (HUB) which is in the implementation period.

The Fat client (IFC) is a desktop solution for the visualization of the results of several workflows. As it can be seen from the figure 4 the functions are going to be embedded and the users can reach each of the components through the main user interface. All the components of the whole system are acting like a black box for the end users but also provide access to certain parameters to serve the expert users needs as well. The client is an internally 3D rendering system with basic navigation mapping functions developed for the specific purposes of the IQmulus project. On the figure 5 the graphical interface of IFC is displayed with one resulting GeoTIFF from the flood and waterlogging detection workflow. The colors of the maps are representing the waterlogging classes like natural water surfaces, waterlog, seriously affected soil, moderately affected soil, weakly affected soil, vegetation in waterlog.

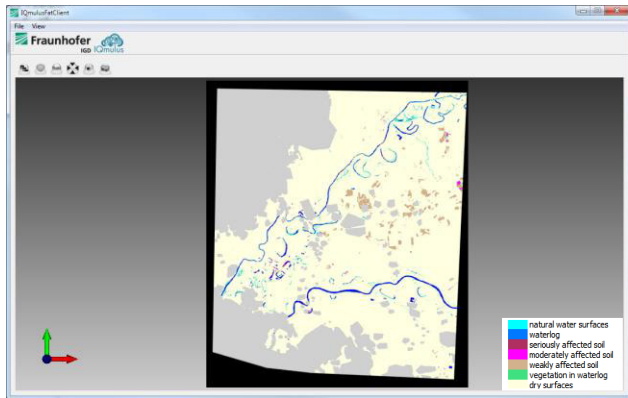


Figure 5: Interface of IQmulus Fat Client (IFC) displaying the GeoTIFF result of the LS3 workflow

4. FLOOD AND WATERLOGGING DETECTION: OPERATIONAL RESULTS

One of the case studies selected for early implementation within IQmulus focuses on the solutions for Rapid Response and Territorial Management. More specifically, we concentrate on the developments and results related to preprocessing and classification of satellite images for the detection of flooded and waterlogged areas. The original, currently used flood detection method developed by FÖMI consists of several steps including preprocessing (geometric transformations, cloud and cloud shadow filtering, radiometric calibration, calculation of spectral indices), feature detection (rule-based classification to provide thematic maps with several categories of water presence). IQmulus plays a major role in increasing the degree of automation of this workflow. Detecting flood and waterlogging is a task that requires a series of satellite images of acquired for different areas, but almost at the same time. Quick response in such a situation is critical, thus the processing time of the satellite and/or aerial image series in a given time is of major importance. Original data needs calibration; other data has to be derived from the original (in this case spectral indices), and the images could be acquired by different sensors.

Based on the above needs, the flood and waterlogging detection workflow has been developed in the frame of the IQmulus project. It is a complex workflow covering all the aspects mentioned above. It consists of multiple algorithms (some of which are also available as separate services). Compared to the current solution, it provides a higher level of automation via smarter algorithms; therefore, it improves overall processing time and implies a better use of human resources.

Radiometric pre-processing (TOA reflectance calculation) and processing of spectral indices is now based on metadata files of satellite imagery, leading to an automatic instead of a manually induces process.

Calculation of spectral indices needed for thematic classification is also based on metadata stored along the images. The whole process can be described and parameterized in Domain Specific Language in the Workflow Editor, and after the selection of datasets to be used it can be launched directly from the IQmulus Graphical User Interface (Figure 6). Results are created in the cloud and can be downloaded for further processing and analysis.

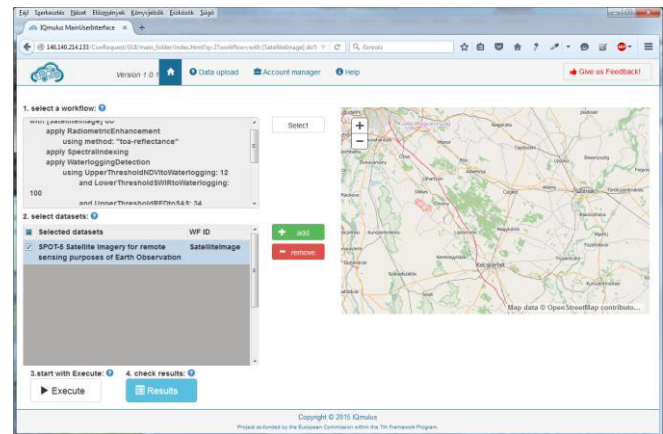


Figure 6: Workflow execution and data selection on the IQmulus graphical user interface

5. ACKNOWLEDGEMENTS

IQmulus (full name: A High-volume Fusion and Analysis Platform for Geospatial Point Clouds, Coverages and Volumetric Data Sets) is a 4-year Integrating Project (IP) partially funded by the European Commission under the Grant Agreement FP7-ICT-2011-318787. It is positioned in the area of Intelligent Information Management within the ICT 2011.4.4 Challenge 4: Technologies for Digital Content and Languages. IQmulus started on November 1, 2012, and will finish on October 31, 2016.

6. REFERENCES

- [1] A. Olasz and B. Nguyen Thai, Decision support on distributed computing environment (IQmulus). OGRS 2014. Proceedings of the 3rd Open Source Geospatial Research & Education Symposium 2014. 06.10-13. pp.107-114.
- [2] M. R. Ghazi and D. Gangodkar, Hadoop, Map/Reduce and HDFS: A Developers Perspective Procedia Computer Science, Volume 48, 2015, pp. 45-50.
- [3] The Apache Software Foundation, "Apache Hadoop," Wiki 2007. Available at <http://wiki.apache.org/hadoop/> [Nov 02, 2015].

[4] G. Wang and Q.Weng, 2014. Remote Sensing of Natural Resources. Taylor and Francis Ltd.,USA. pp 25.

[5] B. Nguyen Thai and A. Olasz, Raster data partitioning for support distributed GIS processing. ISPRS Archives Volume XL-3/W3, 2015 pp.543-551

[6] M. Krämer and I. Senner, A modular software architecture for processing of big geospatial data in the cloud, Computers & Graphics, vol. 49, pp. 69–81, 2015.

[7] D. Kristóf, R. Giachetta, A. Olasz, B. Nguyen Thai, Big geospatial data processing and analysis developments in the IQmulus project; Proceedings of the 2014 Conference on Big Data from Space (BiDS'14) pp. 214-217.

[8] R. Giachetta, AEGIS A framework for processing large scale geospatial and remote sensing data in MapReduce environment ,Computers & Graphics Volume 49, 2015, pp. 37–46

[9] J. R. Herring (ed.) The OpenGIS Implementation Standard for Geographic Information: Simple Feature Access – Common Architecture, version 1.2.1. 2011 Open Geospatial Consortium Inc., OpenGIS Project Document Number 06-103r4, Wayland, MA, USA

[10] N. Golpayegani and M. Halem, “Cloud Computing for Satellite Data Processing on High End Compute Clusters”, IEEE International Conference on Cloud Computing, 2009 (CLOUD '09), IEEE, New York, NY, USA, pp. 88–92

[11] A. Cary, Z. Sun, V. Hristidis and N. Rishe, “Experiences on Processing Spatial Data with MapReduce”, Lecture Notes in Computer Science, Springer-Verlag, Berlin, Heidelberg, Germany, 5566, 2009 pp. 302–319

[12] F.Michel, T. Franke, T. Gierlinger, A. Brodtkorb Interactive visual decision support techniques, Deliverable D5.1.1 IQmulus project documentation 2013.

Improving Effectiveness of Simulation Performance Prediction in Clouds

László Muka and István Derka

Széchenyi István University, Department of Telecommunications, Győr, Hungary

e-mails: muka@sze.hu, steve@sze.hu

Abstract—The cloud computing is becoming increasingly important execution environment for discrete event simulation models too. The cloud services are expected to provide an automatic and effective parallel and distributed execution but the occurring communication delay in the network may decrease the cloud performance. The paper describes closed queuing network model of on-demand resource use in cloud execution. The paper introduces a method based on coupling factor method and on rough set train-and-test approach of effective simulation performance prediction.

Keywords—parallel and distributed simulation, prediction of simulation performance, coupling factor method, cloud execution environment, rough sets, prediction effectiveness

I. INTRODUCTION AND MOTIVATION

Over the last few years, the need for the discrete event simulation (DES) of large-scale complex networks and network services has been constantly growing because DES turned to be an efficient tool for the analysis of these systems. The computing capacity requirements large-scale and complex networks can be fulfilled by parallel and/or distributed discrete event simulation approach [2, 12,13].

According to a common and simple definition [2], Parallel and Distributed Simulation (PADS) is defined as any simulation in which more than one processor is used. PADS is the execution of a single discrete event simulation model on a high performance computing platform: on clusters of homogeneous and heterogeneous computers, on WEB, grid and cloud execution environment.

The typical situations of PADS approach applications from the point of view of the runtime performance requirements: time consuming applications can be for example the simulation of large and/or complex systems and networks.

The development and use of systems based on the PADS methods are resource consuming, not easy task even today. The *simulation performance prediction* method is appropriate for support to reach good PADS performance [11].

The question of support may be formulated in the following way: How to build a model with a good parallelization potential and how to execute it with a good runtime performance involving the necessary (available) resources of a parallel and/or distributed environment and how to do it in an effective way?

The main *contributions* of the paper can be summarized as follows:

- Closed Queuing Network (CQN) Virtual Machine (VM) model of cloud work including modelling of latencies among cloud instances has been defined.
- Closed Queuing Network (CQN) Virtual Machine (VM) model is enhanced to take into account the use of cloud resources (VMs) on on-demand base.
- Rough model of prediction is introduced to support modelling and the performance prediction in an effective way.

The paper is organized as follows. First, the issues of simulation execution in cloud are examined. Then questions of simulation performance are described including the enhanced CQB and VM model of cloud execution. Forth section presents the rough set modelling of performance prediction in an example. The fifth section concludes the work.

II. ISSUES OF SIMULATION EXECUTION IN CLOUD

A. Cloud Services

It has four deployment models private cloud, community cloud, public cloud and hybrid cloud.

Cloud services can be provided according to three service models, e.g., Software as a Service (S-a-a-S), Platform as a Service (P-a-a-S) and Infrastructure as a Service (I-a-a-S) [1].

B. Parallel and Distributed Simulation in the Cloud

Cloud can be the execution environment for PADS with huge capacity requirement since it has made accessible high performance computing platforms (HPC) to end users of the cloud. For example, in the Amazon's very successful Elastic Compute Cloud (EC2), the Message Passing Interface (MPI) – the standard for parallel programming message communications protocol – is supported by the cloud.

Cloud environments are often better at providing high bandwidth communications among applications than in providing *low latency* [15].

PADS applications *typically* work with significant *communication among segments* with sending a lot of short messages between the processes. Thus, for PADS good performance *quick transport* (= *low latency*) is more important than high bandwidth alone

Problems related with functioning of the cloud, for the PADS applications using *optimistic synchronization* protocol, may lead to performance degradations [16].

There has been made researches in cloud computing and but less attention have been paid to parallel and distributed simulation and even less to *conservative synchronization method*.

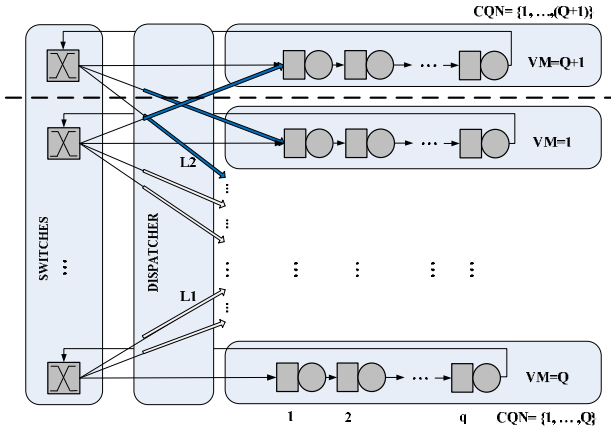


Figure 1. CQN switching and latency model of cloud execution environment (CPU=VM, LP=LP)

III. SIMULATION PERFORMANCE IN CLOUD

A. General Model of Cloud Performance Simulation

Buyya define cloud computing in the following way [21]: “A Cloud is a type of parallel and distributed system consisting of a collection of inter-connected and *virtualized* computers that are *dynamically* provisioned and presented as one or more unified computing resource(s) based on service-level agreements established through negotiation between the service provider and consumers.”

The principle of locality of execution allows to model the PADS cloud executions as an addition of a stable set and a changing set of processes (Figure 1).

The CQN model with tandems (1, ..., Q) of simple queues (q) for the stable set of processes and ((Q+1), (Q+2), ...) for changing set of processes. Logical process of a tandem is assigned to a virtual machine (VM) Modelling of delay in switching between tandems can be used for modelling the PADS performance (represented by L1 and L2 of dispatching in Figure 1). (The CQN model itself is appropriate for modelling of PADS execution [3,4,9,10,14].)

B. Latency and Jitter in Cloud

In this point, the network performance of EC2 will be discussed, focusing on the measurements of packet delay measurement in spatial experiment published in [5].

In the experiments in [5], the packet round trip delay (RTT) has been measured in EC2, for 750 small instance pairs and 150 medium instance pairs using 5000 ping probes.

For the examined instance pairs, the measured hop count values were within 4 hops.

Diagram in Figure 2 (that has been built on data measured in [5]), shows the inverse cumulative distribution function (ICDF) of RTTs for small and medium instance pairs.

The diagram shows, that the RTT values among the examined instances are not stable.

Remarks for the min and max RTT values for the small and medium instance pairs:

- the max RTT values for the 95% of small instances are higher than for the medium instances. Range for small instances is are between 4 and 60 msec
- the min RTT values are higher for medium instances (average difference 0,055 msec)

C. The Coupling Factor Method

The principle of the *Coupling Factor Method (CFM)* of performance prediction [7, 9,17] may be formulated as an inequity:

$$L * E \gg \tau * P$$

where L is the lookahead value characterizing the model (simsec), E is the event density generated by the model (ev/simsec), τ is the latency of messages between logical process (LPs) of the model (sec), and P is the event processing computation hardware performance (ev/sec). According to the method, the coupling factor λ is calculated according to the formulas

$$\lambda = \frac{L * E}{\tau * P}$$

The high value of the coupling factor λ shows the good potential of the simulation model for parallelization. The formula involves only four parameters for the calculation which can be measured in simple sequential simulation runs.

For a separate process, the λ_N parallelization potential of a process is only a part of the whole potential:

$$\lambda_N = \frac{L * E}{\tau * P} * \frac{1}{N_{LP}}$$

where N_{LP} the number of the LPs [8].

D. Setting up the Scale Prediction for a Homogeneous Cluster

For the presented method, the results described in [8] will be used in the cloud performance analysis.

The hardware environment is a *homogeneous cluster* of 12 PCs. For the experiments in the case, the starting number of jobs is 2 jobs/simple queue with exponential inter-arrival time and with exponential service time distributions (the expected value for arrival and service time is 10sec). The delay on links between simple queues is 1sec. The number of tandem queues is 24 (Q), the number of simple queues/tandem queue is 50 (q). The switching in CQN model is performed by uniform distribution to switch to the next tandem queue, and the delay of switching models lookahead.

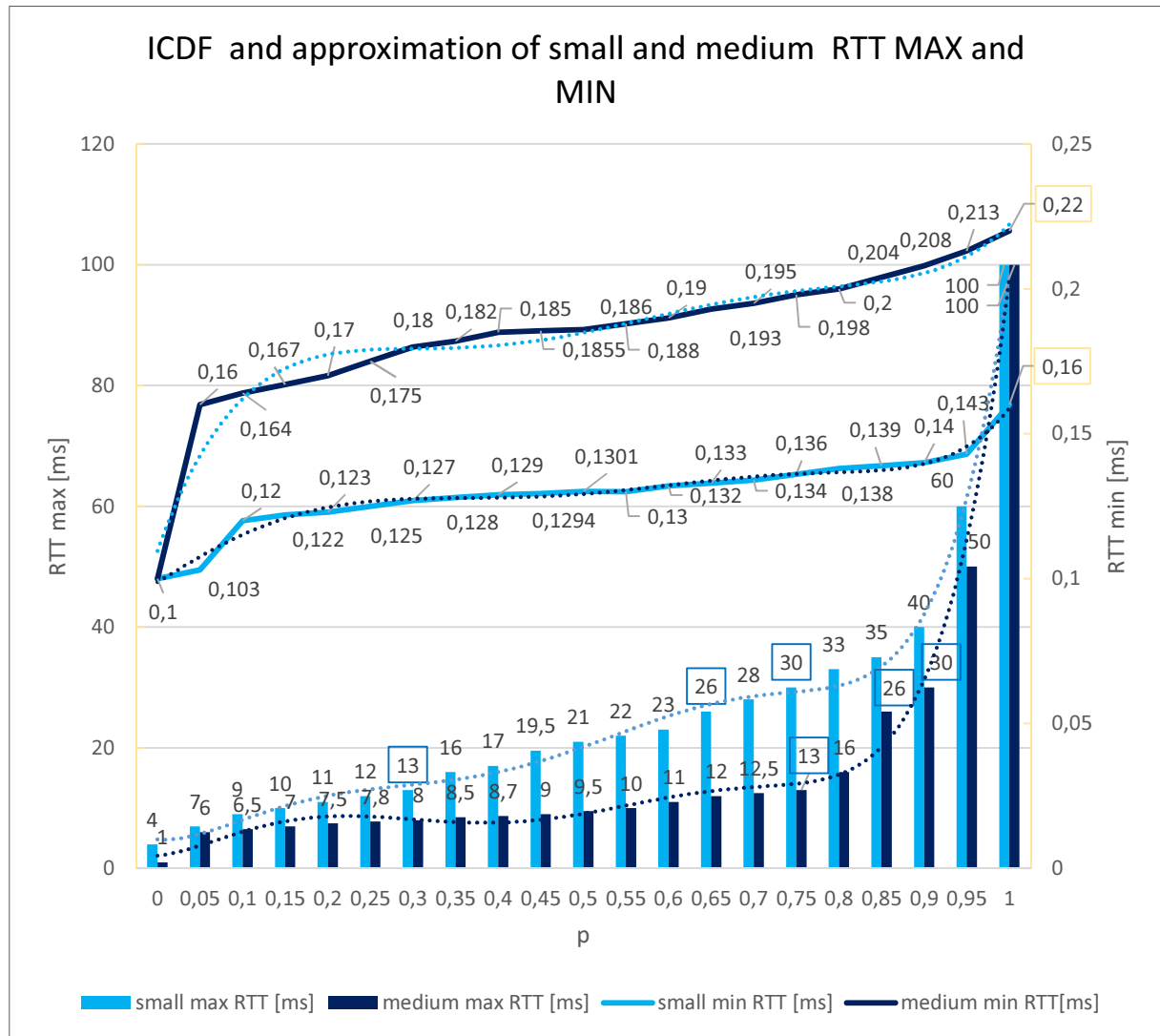


Figure 1. ICDF of max and min RTTs for small and medium sized instances in EC

2

The task assignment principles in the execution are the partitioning into LPs and Load Balancing Criterion.

The measured value of τ is 0,025ms for the system in the case. The value of variables that have been measured in sequential simulation runs are summarized in Table I.

TABLE I.
COUPLING FACTOR MEASUREMENT AND CALCULATION

1.	L	[simsec]	0.1	1	10	100	1000
2.	Number of events	[ev]	138122606	138091806	137816386	134885378	102957082
3.	WCT (N=1)	[sec]	524	521	523	516	416
4.	Simulated virtual time	[simsec]	864000	864000	864000	864000	864000
5.	P	[ev/sec]	263502	264868	263465	261132	247653
6.	E	[ev/simsec]	159	159	159	156	119
7.	τ	[sec]	0.000025	0.000025	0.000025	0.000025	0.000025
8.	R=P/E	[simsec/sec]	1648	1657	1651	1672	2078
9.	L / τ	[simsec/sec]	4000	40000	400000	4000000	40000000
10.	λ measured		2.43	24.1	242	2391	19246

TABLE II.
COUPLING FACTOR MEASUREMENT AND CALCULATION

1.		RTT(ms)	latency(ms)	τ latency	proportion	λ shift (log10)
2.	CQN model τ		0,025			
3.	EC2 average	0,2	0,1	4		0,60
4.	EC2 average	3	0,15	4,2		0,62
5.	EC2 small instances	10	5	200	55%	
6.	EC2 medium instances	20	10	400	20%	
7.	EC2 small instances weighted τ latency		2,84	113,6		2,06
8.	EC2 medium instances weighted τ latency		2,16	86,4		1,94

E. Analysis

Table II summarizes the EC2 latency measurements [5] and the calculated λ shifts in predicted relative speedup [19].

IV. EXAMPLE OF ROUGH SET MODEL OF CLOUD EXECUTION

The *RST model of simulation performance prediction* can be set up in the form of *decision information system (DIS)* and decision tables:

$I = (U, C \cup D, f_{V'}, f, V', V)$, (before discretization and coding)

$I = (U, A = C \cup D, f, V)$,

$I = (U, A = C \cup \{d\}, f, V)$, $d \in D$.

The *simulation experiments* (sequential and PADS runs) will be the *objects* of the universe

$U = \{x_1, x_2, x_3, \dots, x_{|U|}\}$

where x_i denotes the i -th simulation experiment of the universe.

The set of *condition* attributes may be defined as follows

$C_{\text{Cloud CFM model hardware-software features}}$
 $= \{P, E, L1, L2, \lambda, \lambda_N, N_{LP}, \tau(RTT/2), N_{VMS}$

simulated time, runtime (wall – clock time) $\} \cup$

$C_{\text{hardware-software environment}} \{OMNet +$
 $+, MPI, Linux Debian\} \cup$

$C_{\text{Cloud CQN models' features}}$
 $\{conservative synchronisation protocol\} \cup$
 $\{with null message algorithm\}$

$\{Q, q, propagation\ delay\ between\ simple\ queues,$
switching time between tandems,
starting number of jobs of simple queues in cloud CQN models,
inter arrival time of jobs,
service time, service discipline (FCFS) $\}$

In set C , attributes, relating to the coupling factor method, are in accordance with the *homogeneous execution environment* and with the equal (fair load balancing criterion too (implicit) condition attributes)

The set of *decision* attributes under consideration can be set as

$D = \left\{ \begin{array}{l} speedup, R - speedup, \\ number\ of\ vacationing\ jobs\ in\ L1\ and\ L2 \end{array} \right\}$.

The RST exploring of the coupling factor method is performed using the following decision table:

$I_t = (U, A = C \cup \{d_t\}, f, V)$, $d_t \in D$,

The examination is executed in a form of a *train-and-test analysis*:

$U \rightarrow (U_{\text{training}(i)}, U_{\text{test}(i)})(i = 1, 2, 3, \dots, |steps|)$,

$f_{RNDsplit(U,i)} : U \rightarrow (U_{\text{training}}, U_{\text{test}})$ with

$\forall(i)(i = 1, 2, \dots)(U_{\text{training}(i)} \cup U_{\text{test}(i)} = U) \wedge$
 $(|U_{\text{training}(i)}|/|U_{\text{test}(i)}| = \frac{x}{y}) \wedge (RND_{seed} = i)$.

For the *train-and-test* examination, the ROSETTA system [18] can be used with some selected rule generation algorithm G (for example, Johnson's RSES (Rough Set Exploration System)) and with the subsequent classification of objects.

For evaluation purposes, *re-classification rules* $S_{\text{self-training}(i)}$ can be generated for every $U_{\text{test}(i)}$ using the same rule generation and classification algorithm.

For the *efficacy, efficiency and effectiveness analysis*, [20] the simulation *cost* $K[sec]$ expressed in *computing time "consumption"* should also be calculated, including cost of *attributes, rules, simulation experiments and predictions*, for the series of predictions ($K(S_{\text{training}(i)}, i = 1, 2, 3, \dots, |steps| = \text{number of predictions})$)

This analysis provides a feedback to the simulation model and to the the simulation performance model too, supporting model features identification and refinement.

V. CONCLUSION

In the paper, a model of cloud execution has been described which models the network latencies

This is a Closed Queuing Network (CQN) and Virtual Machine (VM) model of cloud work modelling network latencies among cloud instances

In the paper, the CQN and VM model has been extended in to take into account the on-demand base use of cloud resources

In an example, a rough model of prediction is introduced to support effective modelling and performance prediction

REFERENCES

- [1] P. Mell, and T. Grance. The NIST definition of Cloud Computing, National Institute of Standards and Technology (NIST), Special Publication Draft-800-145, 2011. p. 2.
- [2] R. M. Fujimoto, "Parallel Discrete Event Simulation", *Communications of the ACM*, vol. 33, (1990.) no 10, pp. 31-53
- [3] Lencse, G., I. Derka and L. Muka. 2013. "Towards the efficient simulation of telecommunication systems in heterogeneous execution environments". submitted to the *36th International Conference on Telecommunications and Signal Processing (TSP)*, July 2-4, 2013, Rome, Italy.
- [4] R. L. Bagrodia, M. Takai, (2000). "Performance evaluation of conservative algorithms in parallel simulation languages", *Parallel and Distributed Systems, IEEE Transactions on*, 11(4), 395-411.
- [5] Wang, G., and Ng, T. E. (2010, March). The impact of virtualization on network performance of amazon ec2 data center, In *INFOCOM*, 2010 Proceedings IEEE, pp. 1-9.
- [6] Xu, Y., Musgrave, Z., Noble, B., Bailey, M. (2013, April). Bobtail: Avoiding Long Tails in the Cloud. In *NSDI* pp. 329-341.
- [7] A. Varga, Y. A. Sekercioglu and G. K. Egan. "A practical efficiency criterion for the null message algorithm", *Proceedings of the European Simulation Symposium (ESS 2003)*, (Oct. 26-29, 2003, Delft, The Netherlands) SCS International, 81-92.
- [8] G. Lencse and A. Varga, "Performance Prediction of Conservative Parallel Discrete Event Simulation", *Proceedings of the 2010 Industrial Simulation Conference (ISC'2010)* (Budapest, Hungary, 2010. June 7-9.) EUROSIS-ETI, 214-219.
- [9] A. Varga and R. Hornig, "An overview of the OMNeT++ simulation environment", *Proceedings of the 1st international conference on Simulation tools and techniques for communications, networks and systems & workshops*, (Marseille, France, March 3-7, 2008) 1-10.
- [10] A. Varga and A. Y. Sekercioglu, "Parallel Simulation Made Easy with OMNeT++", *Proceedings of the 15th European Simulation*

- Symposium (ESS 2003)*, (Oct. 26-29, 2003, Delft, The Netherlands) SCS International, 493-499
- [11] Z. Juhasz, S. Turner, K. Kuntner, M. Gerzson, "A Performance Analyser and Prediction Tool for Parallel Discrete Event Simulation", *International Journal of Simulation*, vol. 4, no. 1, May 2003. pp. 7-22.
- [12] Kunz, G., Tenbusch, S., Gross, J., Wehrle, K. (2011, July). "Predicting Runtime Performance Bounds of Expanded Parallel Discrete Event Simulations. In *Modeling, Analysis & Simulation of Computer and Telecommunication Systems (MASCOTS), 2011 IEEE 19th International Symposium on* (pp. 359-368). IEEE.
- [13] R. Ewald, A. M. Uhrmacher, "Automating the runtime performance evaluation of simulation algorithms", In: *Simulation Conference (WSC), Proceedings of the 2009 Winter*. IEEE, 2009. pp. 1079-1091.
- [14] Lencse, G. and I. Derka 2013. "Testing the Speed-up of Parallel Discrete Event Simulation in Heterogeneous Execution Environments". submitted to the *2010 Industrial Simulation Conference (ISC'2010)* (Ghent, Belgium, May 22-24, 2013.)
- [15] Malik, A. W., Park, A. J., Fujimoto, R. M. (2010). An optimistic parallel simulation protocol for cloud computing environments", *SCS Modeling and Simulation Magazine*, 1(4).
- [16] Fujimoto, R. M., Malik, A. W., Park, A. (2010). "Parallel and distributed simulation in the cloud" *SCS M&S Magazine*, 3, pp. 1-10.
- [17] Varga, A.; Hornig, R., 2008. "An overview of the OMNeT++ simulation environment", In *Proceedings of the 1st international conference on Simulation tools and techniques for communications, networks and systems & workshops, ICST (Institute for Computer Sciences, Social-Informatics and Telecommunications Engineering)*,
- [18] Komorowski, J.; Øhrn, A.; Skowron, A., 2002. "The ROSETTA Rough Set Software System", In *Handbook of Data Mining and Knowledge Discovery*, W. Klösgen and J. Zytkow (eds.), Oxford University Press, ch. 2-3, ISBN 0-19-511831-6.
- [19] Muka, L.; I Derka, I.; 2014. "Simulation Performance Prediction in Clouds", In: Orosz Gábor Tamás (szerk.), 9th International Symposium on Applied Informatics and Related Areas – AIS2014., Konferencia helye, ideje: Székesfehérvár, Magyarország, 2014.11.12 Székesfehérvár, Óbudai Egyetem, pp. 142-147., ISBN:978-615-5460-21-0
- [20] Muka, L.; I Derka, I.; 2015. "Rough Set Based Efficiency Improvement of Simulation Performance Prediction", In: Orosz Gábor Tamás (szerk.), 10th International Symposium on Applied Informatics and Related Areas – AIS2015, Székesfehérvár, Magyarország, Székesfehérvár, Óbudai Egyetem (accepted for publication)
- [21] Buyya, R., Yeo, C. S.; Venugopal, S., Broberg, J., Brandic, I.; 2009. "Cloud computing and emerging IT platforms: Vision, hype, and reality for delivering computing as the 5th utility", *Future Generation computer systems*, 25(6), pp. 599-616.

BIOGRAPHIES

LÁSZLÓ MUKA graduated in electrical engineering at the Technical University of Lvov in 1976. He got his special engineering degree in digital electronics at the Technical University of Budapest in 1981, and became a university doctor in architectures of CAD systems in 1987. Mr. Muka finished an MBA at Brunel University of London in 1996. Since 1996 he has been working in the area of modeling and simulation of infocommunications systems, including human subsystems. Mr. Muka got his PhD at the Budapest University of Technology and Economics in 2011. Dr. Muka has been working as an Associate Professor for the Széchenyi István University in Győr since the February of 2012. He teaches system simulation and network security

ISTVÁN DERKA received his Msc at Faculty of Electrical Engineering and Informatics at the Technical University of Budapest in 1995. He worked for the Department of Informatics from 1999 to 2003 and since then has been working for Department of Telecommunications, Széchenyi István University in Győr. He teaches Programming of communication systems and Interactive TV systems. He is an Assistant Professor. The area of his research includes multicast routing protocols and IPTV services in large scale networks.