

A Tenant-based Resource Allocation Model Application in a Public Cloud

Wojciech Stolarz, Marek Woda

Tieto Poland Sp. z o.o., Aquarius, Swobodna 1, 50-088 Wroclaw, Poland
Department of Computer Engineering, Wroclaw University of Technology,
Janiszewskiego 11-17, 50-372 Wroclaw, Poland
wojciech.stolarz@tieto.com, marek.woda@pwr.edu.pl

Abstract: The aim of this study is to check the proposed tenant-based resource allocation model in practice. In order to do this, two SaaS systems are developed. The first system utilizes traditional resource scaling based on a number of users. It acts as a reference point for the second system. Conducted tests were focused on measuring over- and underutilization in order to compare cost-effectiveness of the solutions. The tenant-based resource allocation model proved to decrease system's running costs. It also reduces system resource underutilization. Similar research has been done before, but the model was tested only in a private cloud. In this work, the systems are deployed into commercial, public cloud.

Keywords: Cloud computing; multi-tenancy; SaaS; TBRAM

1 Introduction

Cloud computing is gaining more and more interest every year. Cloud computing is not a new technology. It is rather a mixture of technologies existing before, like: grid computing, utility computing, virtualization or autonomic computing [14], and it finds application in other seemingly indirectly related areas, like hybrid wireless networks [16] integration of information systems [17] or high performance simulation [18].

The National Institute of Standards and Technology (NIST)¹ defines cloud computing as “a model for enabling ubiquitous, convenient, on-demand network access to a shared pool of configurable computing resources (e.g., networks, servers, storage, applications, and services) that can be rapidly provisioned and released with minimal management effort or service provider interaction.

¹ NIST Cloud Computing Standards Roadmap, Special Publication 500-291, Version 2, July 2013

This cloud model is composed of five essential characteristics, three service models, and four deployment models”.

This approach allows Internet based applications to work in distributed and virtualized cloud environment. It is characterized by on-demand resources and pay-per-use [1] pricing. Nowadays, every respected IT-company has started to think about providing its services in the cloud [5]. Currently Cloud computing is one of major enablers for the manufacturing industry [3]. It became widely used to enhance many other aspects of industrial commerce by moving business processes to the cloud to improve the companies’ operational efficiency. Currently, a new trend can be observed, inspired by cloud computing, it is illustrated in the movement from production-oriented manufacturing to service-oriented manufacturing. It converges networked manufacturing, manufacturing grid (MGrid), virtual manufacturing, agile manufacturing and Internet of Things into cloud manufacturing [2], [3], where distributed resources provided by cloud services are managed in a centralized way. Clients can use the cloud services according to their requirements. Cloud users can request services ranging from product design, manufacturing, testing, management and all other stages of a product life cycle [3].

Software-as-a-Service (SaaS) is software delivery model in which entire application (software and data) is hosted in one place (usually in the cloud). The SaaS application is typically accessed by the users via a web browser. It is the top layer in cloud computing stack. SaaS evolved from *SOA (Service Oriented Architecture)* and manages applications running in the cloud. It is also seen as a model that extends the idea of *Application Service Providers (ASP)*. *ASP* is primary centralized computing model from the 1990s. SaaS platform can be characterized by: service provider development, Internet accessibility, off-premises hosting and pay-as-you-go pricing [3]. The SaaS platform supports hosting for many application providers. As opposed to *ASP* model, SaaS provides fine-grained usage monitoring and metrics [8]. It allows tenants to pay according to the usage of specific cloud resources. SaaS applications often conform to multi-tenant architecture, which allows a single instance of a program to be used by many tenants (subscribers) [3]. This architecture also helps to serve more users because of a more efficient resource management than in multiple instances approach [4].

In spite of the fact that the idea of cloud computing utilization has become a reality, questions like how to enhance resource utilization and reduce the resource and energy consumption are still not effectively addressed [1].

Since most cloud services providers charge for the resource use, it is important to create resource efficient applications. One of the ways to achieve this is multi-tenant architecture of SaaS applications. It allows the application for efficient self-management of the available resources.

Despite the fact, that in the cloud, one can automatically receive on-demand resources, one can still encounter problems related to inappropriate utilization of the available resource pool at a particular time. These issues manifest in over- and underutilization, which exists, because of the “not fully elastic”, pay-per-use model used nowadays [10]. Over provisioning arises when, after receiving additional resources (in reply for peak loads), one keeps them, even if they are not needed any more. Such a situation is called underutilization. Under provisioning (saturation) arises when one cannot deliver required level of service because of insufficient performance. This is also known as an overutilization. It leads to customers turnover and revenue losses [1]. *Amazon Elastic Cloud Computing (EC2)* service charges users for every partial hour they reserve each EC2 node. Paying for server-hours is common among cloud providers. That is why it is very important to utilize fully given resources in order to really pay just for what we use.

Due to the fact that we are still in early stages of cloud computing development, we cannot expect cost-effective pay-per-use model for SaaS applications after just deploying it in the cloud. What is more, automatic cloud scalability will not work efficiently if applications consume resources, which are indispensable to meet the desired performance levels [13]. To achieve desired scalability we need to design our SaaS application with that in mind. In order to do so, the application must be aware how it is used [7]. We can use multi-tenant architecture [9] to manage the application behavior. It allows using a single instance of the program for many users. It works in similar way like a singleton class in object programming languages, which can supervise creation and life cycle of objects derived from that class. Supporting multiple users is a very important design step for SaaS applications [2]. We can distinguish two kinds of multi-tenancy patterns: multiple instances (every tenant has got their own instance running on shared resources) and native multi-tenancy (single instance running on distributed resources)² [2]. The first pattern scales well for a small number of users, but if there are more than hundreds of users, it is better to use the second pattern.

The one of most recent solutions for over- and under- utilization problems may be a tenant-based resource allocation model (TBRAM) for SaaS applications. That solution was introduced and tested with regard to CPU and memory utilization by various authors [3]. They proved the validity of TBRAM through the reduction of used server-hours as well as improving the resources utilization. However, the authors deployed their solution into a private cloud which can only mimic a public cloud environment. They tested cases with incremental and peak workload of the system. In this paper we wanted to check whether the TBRAM is really a valuable system. Examining the TBRAM system in a public and commercial cloud environment could deliver the answer to that question. Therefore, the main aim of

² Architecture Strategies for Catching the Long Tail: 2006.
<http://msdn.microsoft.com/en-us/library/aa479069.aspx>. Accessed: 2014-09-07

the paper is the further validation of TBRAM approach, as it was proposed in the future research part of the base work [3]. If the results of the study confirm improvement of the model performance, then it could be considered as the solution to the previously mentioned provisioning problems.

2 System Design

2.1 Base System

In this section the base tenant-unaware resource allocation SaaS system (Base System) is described. It conforms to a traditional approach to scaling resources in a cloud and is based on the number of users of the system. It is substituted by *Elastic Load Balancer* service. According to AWS Developer Forum³ the *Elastic Load Balancer (ELB)* sends special requests to balancing domain's instances to check their statuses (health check) it then round-robins among the healthy instances, having fewer outstanding requests pending. Although the name of this system suggests lack of awareness of tenants it concerns only resource allocation. The system was built according to *Service Oriented Architecture (SOA)* and native multi-tenancy pattern. First, it was implemented as a set of J2EE web applications using *Spring* and *Struts* frameworks. Several parts of the system were later transformed to web services using *WSO2* and *Axis2*. Deploying application as a web service makes it independent of the running platform. It also gives more flexibility when accessing the application.

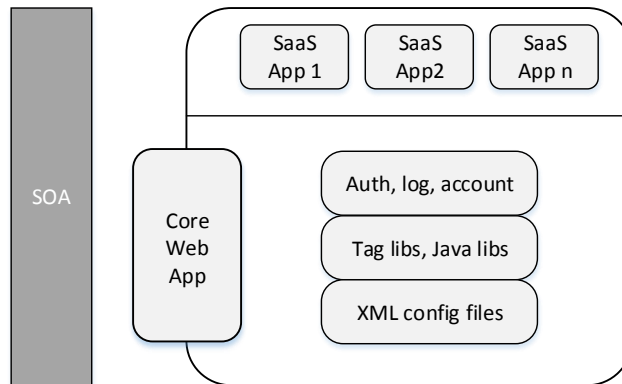


Figure 1

SaaS platform system architecture

³ <https://forums.aws.amazon.com/thread.jspa?messageID=135549&>. Accessed: 2014-10-17

A general, high level, overview of the test-bed architecture is shown in Fig. 3, one can see there a group of *Amazon EC2 (Tomcat)* instances. The number of the instances varies and it depends on the number of simulated users. Each of the instances consists of a Virtual Machine (VM) with one Tomcat web container. In each Tomcat container authorial a SaaS platform is deployed. The platform is the main part of the system as it makes a basis for SaaS applications. It also includes web services and common libraries.

The SaaS platform (Fig. 1) was the main part of the system. The task of the platform was to support deployment of plain web applications as a SaaS service. The idea behind the design of this part was inspired by this work [3]. Their SaaS platform was developed as a part of the *Rapid Product Realization for Developing Markets Using Emerging Technologies* research at *Tecnológico de Monterrey University, Mexico*.

Since the authors had limited means and limited time, The SaaS platform was simplified and focused only on the usability for SaaS system.

The entry point to the platform was the *Core Web App (CWA)*. As the name suggests, it was a web application that worked as a gate. All interactions between outside environment and the parts of the platform were done through this element. In the background, there were applications responsible for users' authorization, account management and logging. The platform contained also common Java libraries used by deployed applications. Configuration was made by the XML or plain text files. The platform exhibited web service interfaces to be consumed by outside applications. One example of such an interface was the interface for metering services. It allows monitoring the usage of resources by the platform. That behavior is depicted by the SOA element in Fig. 1. On top of this, we can see the SaaS applications. These were developed as normal Java web applications, but when deployed on the platform, they gain extra SaaS functionality. Two were implemented, a Sales application and a Contacts Application. It is assumed that two applications are enough to present the platform's functionality as well as the interactions between the deployed applications. It is worth mentioning that there is one more feature, which was not depicted in Fig. 1, the communication channel between the platform and an external database server.

In order to measure over- and under- utilization certain metrics from running VMs were needed. Some of them were available directly through *Amazon CloudWatch* metering service (CPU usage, network in/out and number of requests per second in case of the ELB). The latter metric can be used to calculate the throughput of the entire system. Other source of data for this metric came from *JMeter* tool. However, monitoring of RAM consumption and number of running threads was not provided by the *CloudWatch*. That is why authors developed a monitoring solution called *Resource Consumption Monitor (RCM)* – a service that sat between the monitoring domain and the *CloudWatch*.

There are two main approaches to the monitoring problem. The first is a distributed approach. It is similar to *Observer Pattern* [6] known from object oriented programming patterns and best practices. In this case monitored VMs register themselves to the monitoring service and then publish their measurements. The main strength here, is build simplicity. It has, however, one big disadvantage – each worker VM needs to be aware of the monitoring process. It is also hard to quickly notice a VM termination due to unexpected events or errors. That is why the second, centralized approach was chosen. In this case, the VMs with SaaS platforms are unaware of being monitored as it is beyond their consideration. The RCM is constantly monitoring the state of VMs by polling AWS cloud (the performance hit was negligible, it was not included in the considerations). After each interval (Polling interval) it collects the measurements from the monitored domain. After another interval (Publishing interval) it publishes collected data to the *CloudWatch* service. Thanks to that, all of the VMs measurements were available in one place.

The *RCM* is a web application, but it can also be used as a standalone console Java application packed into an archive file (jar). It used the *Java Management Extension (JMX)* RMI-based protocol which allows to request information about running Java Virtual Machine. Generally, it is recommended to use an authorial web services to fulfil the same tasks, but since the entire SaaS system is too overly complex, the flexibility offered by web services seems not to be really needed. Especially, that such flexibility comes with a price. First of all, the JMX packets are much smaller than competitive SOAP protocol ones. Therefore, it reduces network traffic and time necessary to decode the packet. The next reason is the requirement for management of web services like Axis2. The *JMX* is built in *Java Runtime Environment (JRE 1.7)*, which is used by authors. Finally, the *JMX* technology is far more robust and advanced. It would be difficult to build a better web service within such a short time frame. It is also transparent for applications running on *JVM*. All what is necessary to do, is to add extra running parameters when starting the *JVM*.

The *RCM* requires a set of parameters to run. One of them is the running mode which tells the monitor whether to run in test mode (very frequent data collecting, but without publishing them to the *CloudWatch*) or in normal mode (with synchronization to the *CloudWatch*). The chosen mode affected both (polling and publishing) intervals. In test mode the data were gathered every 5 seconds. In normal mode polling was set to 10 seconds and publishing to 1 minute. These settings matched the settings of *CloudWatch* service working in detailed mode. Using *RCM* one could also manually start or stop the monitoring of certain VM. To tell the *RCM* which instances to monitor a special tag to VMs was added. The tagging is a feature in AWS cloud that helps to organize running instances. The most common usage of tags is for giving names to the instances which are often more meaningful than their IDs.

Because owned metrics are sent to the *CloudWatch* it was crucial that all the measurements (direct and own) for given VM are taken in the same time. They could be published asynchronously, because they contained a timestamp tag. However, the measurement data itself needed to be synchronized in time. Otherwise they would be invalid. To avoid that synchronization mechanism was implemented in the RCM. It was assured that own measurements data are collected at the same moment as the direct *CloudWatch* data.

The RCM was deployed on a dedicated *t1.micro EC2* instance, because it *should* not affect the work of virtual machines it monitored. Thanks to its web interface, it can be managed from any computer via a web browser.

2.2 Tenant Aware System

The base SaaS system was implemented as a reference tenant-unaware resource allocation system. The main flaw of its design was rigid management of VM instances in the cloud. Thus, it could lead to serious over- and underutilization problems. In Chapter III, we show this based on test results. One of the ways to tackle aforementioned issues was proposed in [3] as a tenant-based resource allocation model (TBRAM) for scaling SaaS applications over a cloud infrastructure. By minimizing utilization problems, it should decrease the final cost of running the system in the cloud.

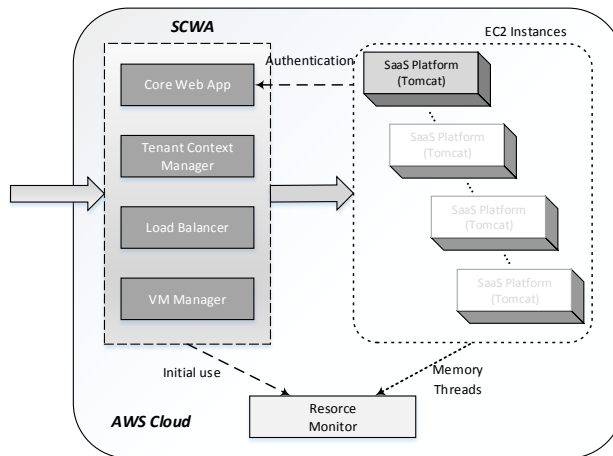


Figure 2

TBRAM system architecture

The TBRAM consists of three approaches that use multi-tenancy to achieve its goals. The first approach is tenant-based isolation [11], [12], which separates contexts for different tenants. It was implemented with tenant-based authentication and data persistence as a part of the SaaS platform (Tomcat instances). The second, is to use tenant-based VM allocation [11], [12]. With this

approach we are able to calculate the actual number of the VMs needed by each tenant in a given moment. The third approach is the tenant-based load balancer that allows to distribute the virtual machines' load with respect to certain tenant. An overview of the architecture is presented in Fig. 2. The dashed line in the picture denotes communication with web services. One can notice that the *SaaS Core Web App (SCWA)* element in the Fig. 2 is the only change made to the original test-bed [3].

A simple load balancer based on round-robin IP address algorithm, is not the best solution to isolate each tenant. Since users from the same tenant share certain tenant-related data it would be a good idea to dispatch their requests to the same VM (if possible). That could reduce the amount of tenant data kept by each VM since some of them would be serving only a few tenants. That could also lead to better usage of servers' cache mechanisms by concentrating on data that are really shared by number of tenant's users. That in turn could for example, reduce a number of requests to database engine.

The key task of the load balancer was to isolate requests from different tenants. The tenant-based load balancer worked in 7th layer of the OSI model. It used information stored in the session context as well as local applications data to assign the load efficiently. The idea behind request scheduling is that requests from one tenant should be processed on the same VMs. If that is impossible, then the number requests should be limited, so they were not scattered along the whole balancing domain. That can also allow reducing context switching and using previously cached data. The traditional scheduling process uses only current status data, so it does not belong to dynamic load balancers family, but TBRAM-based solution is based on adaptive models of load balancing.

As suggested in other work [3], we made the load balancer a part of *SaaS Core Web App (SCWA)*. It was a natural choice to put that element there, since all the requests came through it anyway (because of the centralized authorization service). Design of the load balancer is similar to the one proposed in current work [3]. It consisted of five elements: *Request Processor*, *Server Preparer*, *Cookie Manager*, *Response Parser* and *Tenant Request Scheduler*. Each of them was responsible for specific function in the processing pipeline sequence. The most important was the last part of processing assigned to *Tenant Request Scheduler*. The scheduling policy enforces that the subsequent requests from the same tenant should be dispatched to the same VM. If a given VM was saturated, then the scheduler dispatched the request to the next available VM of that tenant. Finally, if no other VM was available, the scheduler requested a new VM from the VM Manager.

The HTTP as the Internet protocol was designed to be stateless. It means that every request is independent. It starts from handshaking in order to establish a connection. Then data exchange appears for one or possibly more server's resources. After that the connection is closed. When user requests another

resource the whole procedure repeats. However there was a need to keep a track of users action for example to make functioning of shopping chart possible in online shops. Because of private IP addresses it was not feasible to recognize all users just based on their IP. This is where the session mechanism comes with help. In general, it allows storing user related data at the server side and therefore distinguish each unique user. It works fine when there is only one server dealing with a given user, because of the limited session scope. If there are more servers this has to be handled differently. One of the solutions for that problem is clustering of Tomcat servers. But even better solution is to dispatch given user's requests in a unique server as it eliminates the need of session sharing. For that purpose many available load balancers offer so called session stickiness or session affinity. This feature allows grouping requests of a given user within a session scope and sending these requests to the same server. When it comes to tenant-based load balancer it could be called tenant stickiness or affinity. It can be imagined as yet another layer above the session layer which groups requests from a given tenant.

3 Test Results and Analysis

We tested our TBRAM-based SaaS system and compared it with the non-TBRAM version. The comparison was made in terms of overutilization, underutilization and financial cost. To calculate cost of running certain SaaS system in the cloud, billing statement delivered by Amazon was used. To measure over- and under-utilization of resources, measurements data collected by *Amazon CloudWatch* monitoring service were used. Before this was done the entire test bed was deployed in Amazon cloud environment (AWS). Then the system was stressed with the workload of HTTP requests. We used a cluster of *JMeter* machines performing test plans to achieve that. There was one test bed. The main difference in the architecture was in the entry point to the SaaS system. In the case of TBRAM it was the SCWA element including a load balancer, tenant context manager and virtual machines manager. In the case of the standard model it was just the *Elastic Load Balancer (ELB)* service from Amazon.

3.1 Over- and Under- utilization Results

Test results were collected by Amazon CloudWatch monitoring service. This tool allows to view some basic statistics of data in form of charts. However, in order to perform more advance analysis it was needed to download the raw data for further processing. The results are presented in following tables (Table 1, Table 2).

Table 1 presents results of the tests performed over the Base System that used traditional resource scaling. The results come from memory and CPU monitoring of the system. The first column contains the months of simulated year (24 hours of

tests). The second column presents the number of virtual machines running in each simulated month.

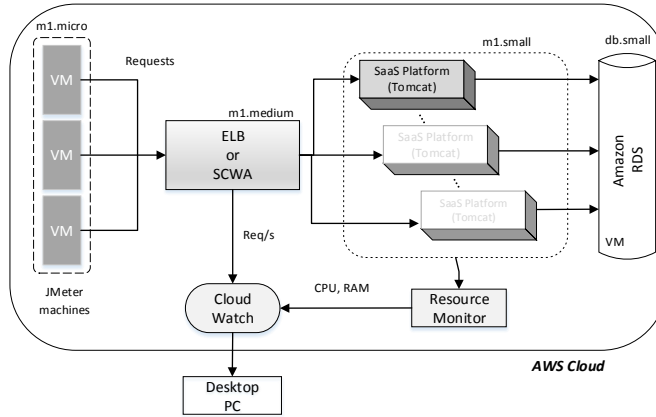


Figure 3

Test bed architecture

This number is valid only for the incremental workload tests. The peak-based tests were conducted in slightly different way. Instead of time constraints they were set to perform certain number of test plan iterations. We can notice the difference between the test types in server hours provided by the VMs each simulated month.

Table 1

Results of the Base System tests

Simulated month	VMs	Server-hours		Combined-incremental		Combined-peak	
		incremental	Peak-based	UU (%)	OU (%)	UU (%)	OU (%)
January	2	1440	2460	38.75	0.00	31.94	0.00
February	2	1440	2580	20.83	0.83	34.17	0.00
March	2	1440	2460	0.00	0.00	35.42	0.00
April	2	1440	2580	0.00	9.15	37.22	0.00
May	4	2880	4200	19.38	6.66	43.36	0.98
June	4	2880	5160	9.79	0.00	26.39	0.00
July	4	2880	4920	10.63	0.00	19.03	0.00
August	4	2880	8040	10.83	0.00	64.31	0.49
September	8	5760	9840	31.61	0.00	7.43	0.49
October	8	5760	10320	34.48	0.00	0.21	1.46
November	8	5760	9840	19.90	0.00	1.56	0.49
December	8	5760	7440	21.39	0.00	56.04	0.00
Total		40320	69840				
			Avg.	18.13	1.39	29.76	0.33

In the case of the incremental test the total value can be simply calculated by multiplying the number of VMs by the number of hours in the month (number of VMs * 24h * 30 days). In the case of the peak-based test such calculation is not straightforward. This is because peak-based tests were little longer than the original time frame of 24 hours (simulated year) per each test. The last four columns of the table contain *combined utilization*. This term describes a situation during the tests when a given VM was saturated or underutilized with respect to the both measures (CPU and memory). The combined utilization percentages were calculated based on formulas [3]:

$$\%UU(\text{Underutilization}) = \frac{\left(\frac{\text{Combined } UU}{\text{Measurements per hour}} \right)}{\text{server hours}} \quad (1)$$

Formula (1) calculates the combined underutilization of a given VM per each time period. It yields a percentage of wasted VMs out of all available in that time. We could also divide the number of measurements when a VM was underutilized by the number of all measurements to get the information how often the UU occurred. This number oscillated around 50% for both test types. However, the measure defined by the formula (1) is more informative since it shows the size of the problem, not just the occurrence frequency. In case of overutilization the following formula was used:

$$\%OU(\text{Overutilization}) = \frac{\text{Combined } OU}{(\text{Measurements per month} * \text{VMs number})} \quad (2)$$

Overutilization informs us about a percentage of VMs that were saturated each month. It is calculated by dividing the number of measurements that had inflection points by all measurements performed during the measured time period (measurements per month * VMs number). As it can be noticed OU hardly appear in the tests. This is because the VMs workload was chosen with overutilization in mind. We did not want to saturate the VMs too much, but during the final tests the system behaved even better than expected. Therefore saturation of the machines was lower. Nevertheless, the tests of TBRAM system were conducted under exactly the same conditions, so it shouldn't bias the results. One can also observe the total number of server-hours provided by the SaaS platform VMs. It was 40320 and 69840 for the incremental and peak-based tests respectively.

Table 2 shows results of the tests of the TBRAM system which uses tenant-based load balancing and resource scaling. Since the VMs fleet size was adjusted to the current needs dynamically there is no corresponding VMs number column with fixed size for each month. The first observation is that the total server-hours were reduced by 19.94% and 30.21%, for the incremental and peak-based tests, respectively. The %OU was marginally smaller as in the case of the Base System. There was however a difference in %UU between the systems. First of all, we can notice that underutilization for incremental workload was smaller at the beginning of the simulated year as compared with the Base System. This is at least partially caused by the dynamic scaling method of TBRAM system. Whereas, the Base

System started with 2 VMs the second system could increase this number starting from only one machine. As it can be seen in Table 2 TBRAM system did not perform so well in the second part of the year.

Table 2
Results of TBRAM system tests

Simulated month	Server-hours		Combined-incremental		Combined-peak	
	incremental	Peak-based	UU (%)	OU (%)	UU (%)	OU (%)
January	720	1200	0.00	0.00	0.00	0.00
February	768	2040	0.00	0.81	13.52	0.65
March	1440	1920	0.00	2.44	0.83	0.65
April	1440	1440	0.00	1.63	26.25	3.26
May	1800	4440	2.78	0.00	0.00	0.00
June	2160	3960	10.56	6.50	2.08	2.60
July	2880	2280	15.83	0.81	16.89	2.60
August	3240	6000	1.19	0.00	31.54	0.65
September	3600	7680	29.17	0.81	1.62	0.65
October	3960	6720	40.02	1.63	20.94	0.00
November	4680	6000	54.45	0.81	25.00	0.65
December	5586	5064	51.59	1.63	35.42	0.00
Total	32274	48744				
		Avg.	17.13	1.42	14.51	0.98

It is important to notice that both averages for %UU are generally lower than in the case of traditional scaling system. In order to check if the TBRAM system leads to the significant improvement over the Base system, we used the t-test to compare UU% in peak-based test:

Table 3
Parameters of the t-student test

N	Degrees	Accuracy	α	t_α
12	22	97.50%	0.025	2.07

where, N is the number of samples (months). Degrees stands for degrees of freedom of the t-test and it is equal to $(N_1 + N_2 - 2 = 12 + 12 - 2 = 22)$. Unlike the authors of the base paper the accuracy was chosen to be set to 97.5% rather than 99.5% because it was thought that test conducted in real public cloud environment is less predictable than a private cloud cluster. The t_α is a base parameter value from the t-student distribution table for the significance $\alpha = 0.025$. If the t-student value for given columns in both tables is greater than the base parameter (2.074) then one can say with 97.5% certainty, that the column's averages are significantly different. We can see that the t-student value for the %UU in peak-based test is equal to 2.1854 (>2.074). Therefore the TBRAM system statistically improved that characteristic. However, %OU averages were

not improved according to the t-student test. The t-student values are given in TABLE3. They were calculated using the following formula:

$$t = \frac{x_1 - x_2}{\sqrt{\left(\frac{s_1^2}{N_1} + \frac{s_2^2}{N_2}\right)}} \quad (3)$$

where x_1 is an average and s_1 is a standard deviation of a column in Table1, where x_2 and s_2 are respective values for a column in Table 2.

3.2 Cost Analysis

Before we explain the cost analysis, a brief description of AWS pricing model is needed. Even despite this vendor specific model, the general idea of pay-per-use is common with cloud computing providers.

One of the main reasons for using the cloud infrastructure is its flexibility. AWS model is also flexible and is based on either pay-as-you-go and pay-per-use. The first one means there is no need for long term contracts neither for any minimal commitment. The latter one is strongly related to utility computing roots of the cloud computing. It means that we pay for what we used. Abovementioned fact is generally true, but with certain granularity, for example: each started running hour of EC2 instance. In case of the AWS there is also no need to pay up-front for any resource. We are also free to over-utilize or under-utilize our resources without any additional fees. There are three fundamental characteristic for which one pays in the AWS. These are: *CPU*, *Storage* and *Transfer OUT*.

In case of *computing*, we pay for each partial hour of our resources from start to stop of the instance. If it comes to *storage* we pay per each GB of stored data. There are many usage ranges with different prices. The more data we store the less per GB we pay. The *transfer OUT* is generally considered as data transferred out of the AWS resources through the Internet. Transfers between our AWS resources do not count in the same way. Communication between the resources within the same *Availability Zone* (distinct physical location belonging to certain region) is free of charge. What is also important is that we do not pay for any inbound traffic to our cloud resources. It does not matter whether the *transfer IN* comes from our other resources or from the Internet.

Table 4
EC2 SaaS platform instances cost comparison

	Incremental	Peak-based	Total
Base System	10.20 USD	16.49 USD	26.69 USD
TBRAM	7.57 USD	12.24 USD	19.81 USD
Total	17.77 USD	28.73 USD	46.50 USD

We used the *AWS* part to compare our system with the *EC2* computing service. This included the instances running the SaaS platform, *ELB* and *CloudWatch* monitoring. It also included the *AutoScaling*⁴ service, but it was free to use as a tool. One pays just for the outputs of that tool, like increased number of running *EC2* instances. Except the server-hours mentioned before, the *EC2* cost depends on: instances type (m1.small, m1.medium and *ELB* on-demand instances in this case) and of course the number of instances. It is worth mentioning, that *Amazon* charges additionally for the amount of data processed by *ELB* load balancer. The last *EC2* service that was included in the cost calculation was running *CloudWatch* in detailed monitoring mode. We were charged only one time per use for both tested systems. Therefore, it was excluded from the cost comparison.

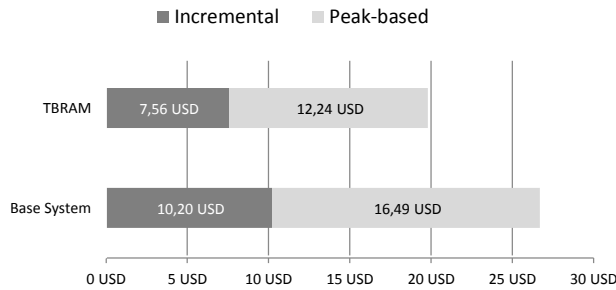


Figure 4

The SaaS platforms costs comparison

Table 4 shows the cost of both systems excluding the load-balancing cost (*ELB* or *SCWA*). That means that only costs of running the SaaS platform VMs were included. Fig. 4 we visualize the costs distribution for each test (simulated year = 24hour of real tests). The TB RAM cost is again, lower than the Base System's with 25.8% improvement. This holds even in case of incremental workload test when the TB RAM system did not statistically improve the underutilization.

Table 5

EC2 load balancing cost comparison

	Incremental		Peak-based		Total	
	cost	%system cost	cost	%system cost	cost	%system cost
ELB	3.20 USD	23.88%	5.25 USD	24.14%	8.45 USD	24.04%
m1.medium	4.25 USD	35.97%	5.61 USD	31.43%	9.86 USD	33.24%
Total	7.45 USD	29.55%	10.86 USD	27.43%	18.31 USD	28.25%

⁴ Amazon Auto Scaling: <http://aws.amazon.com/autoscaling/>

However, the biggest cost reduction is for the peak-based workload tests. It needs to be remembered that the SCWA contained also other parts of the system. Therefore the TBRAM system could not possibly work without that component.

The next step was to compare the separated costs of load-balancing with the results presented in Table 5. The base system used ELB as a load balancer and the authorial TBRAM system used *m1.medium* instance with the SCWA (that contained load balancer). The column named *%system cost* shows what part of the total system cost constitute the load-balancing part. The ELB made 24.04% of the Base System cost, when the *m1.medium* instance made 33.24% of the TBRAM system cost. The ELB was also 14.2% cheaper in terms of USD price. This was because the cost of *m1.medium* instance per hour is over 6 times more than the ELB. We can see that, even despite additional data processing and transfer cost that the ELB introduces, the Amazon's load-balancing service was cheaper. The main difference was in the ELBs scaling up ability. One could notice the moment when the ELB scaled up during the preliminary tests. This motivated us to build up the authorial load balancer deployed into more powerful EC2 instance than the standard *m1.small*. That was clearly an over-provisioning for the time when the system load was low. The lack of enterprise scalability of the SCWA (for small (<10 tenants) scale implementation is equally good as ELB) was the main reason for the higher load-balancing cost of the TBRAM system. It is valuable to notice, that in spite of the fact, that ELB use is cheaper it is not so "resource wise" as the proposed load balancer used in SCWA approach.

Conclusions

This work was inspired and based on the base paper [3]. We wanted to check in practice if the model proposed in the base paper can really influence cost-effectiveness of SaaS systems running in a public cloud. As opposed to just testing the model in the private Eucalyptus cloud. Comparing results from [3] with ours, great similarities were shown. In the base paper [3] the authors achieved 32% server-hours reduction compared to traditional resource scaling. In this work we achieved about 20% and 30% reduction in case of incremental and peak-based tests, respectively. Better result for the peak-based test is caused mainly by the TBRAM underutilization improvement achieved for this type of workload. In the base work the model statistically improved also only the underutilization, but for both types of workload. Thus, we think that this work confirms the TBRAM benefits making it worthwhile, in practice.

Development and deployment of the SaaS systems into AWS cloud made us to draw some other conclusions, too. First, this research showed that TBRAM can improve cost-effectiveness. On the other hand, conformance to that model introduces non-negligible development overhead. This is because we need to write the code for the proposed load balancer, a VM manager (scaling) and a resource monitor. These are not trivial elements to implement and have a great influence on the overall system performance. They are also not easy to test. Because most of

the system's components are independent and distributed, practically the only place they can be fully tested is within the cloud environment, in which, they are implemented. Thus, they make our system more complex and error-prone. Without using the TBram one could simply utilize the robust and flexible services delivered by a cloud provider like *Elastic Load Balancer* and *CloudWatch* monitoring for the case of Amazon, which are not so resource efficient, but noticeably cheaper. So, the model introduces additional costs for the system development and deployment. It is up to us to calculate whether the one-time cost will be returned by the possible savings from a decreased, system running costs.

References

- [1] M. Armbrust, et al.: Above the Clouds - A Berkeley View of Cloud Computing. EECS Department University of California, Berkeley Technical Report No. UCB/EECS-2009-28 February 10, 2009
- [2] Jie Guo Chang et al. (2007) A Framework for Native Multi-Tenancy Application Development and Management. 2007 9th IEEE International Conference on e-Commerce Technology and the 4th IEEE International Conference on Enterprise Computing, e-Commerce, and e-Services, 23-26 July 2007 (Piscataway, NJ, USA, 2007) pp. 470-477
- [3] J. Espadas, A. Molina, G. Jimenez, M. Molina, R. Ramirez, and D. Concha: A Tenant-based Resource Allocation Model for Scaling Software-as-a-Service Applications over Cloud Computing Infrastructures, Future Generation Computer Systems Vol. 29, Issue 1, January 2013, pp. 273-286
- [4] C. Hong et al.: An end-to-end Methodology and Toolkit for Fine Granularity SaaS-ization. 2009 IEEE International Conference on Cloud Computing (CLOUD) 21-25 Sept. 2009 (Piscataway, NJ, USA) pp. 101-108
- [5] R. Iyer, et al. VM3: Measuring, Modeling and Managing VM-shared Resources. Computer Networks. 53, 17 (Dec. 2009) pp. 2873-2887
- [6] R. C. Martin, M. Martin: Agile Principles, Patterns, and Practices in C#. 2006, Prentice Hall
- [7] G. V. Mc Evoy, B. Schulze (2008) Using Clouds to Address Grid Limitations. 6th International Workshop on Middleware for Grid Computing, MGC'08, held at the ACM/IFIP/USENIX 9th International Middleware Conference, December 1-5, 2008
- [8] X. Meng, et al. (2010) Efficient Resource Provisioning in Compute Clouds via VM Multiplexing. 7th IEEE/ACM International Conference on Autonomic Computing and Communications, ICAC-2010 and Co-located Workshops, June 7-11, 2010 (Washington, 2010) pp. 11-20

- [9] M. Pathirage, A Multi-Tenant Architecture for Business Process Executions 2011 IEEE International Conference on Web Services (ICWS) pp. 121-128 ISBN 978-1-4577-0842-8
- [10] M. Stillwell, et al. Resource Allocation Algorithms for Virtualized Service Hosting Platforms. *Journal of Parallel and Distributed Computing*. 70, 9 2010, pp. 962-974
- [11] W. Stolarz, M. Woda: Proposal of Cost-Effective Tenant-based Resource Allocation Model for a SaaS System. *New Results in Dependability and Computer Systems: Proceedings of the 8th International Conference on Dependability and Complex Systems DepCoS-RELCOMEX*, September 9-13, 2013, Brunów, Poland / Wojciech Zamojski [et al.] (eds.) Springer, 2013, pp. 409-420
- [12] W. Stolarz, M. Woda: Performance Aspect of SaaS Application Based on Tenant-based Allocation Model in a Public Cloud. *Proceedings of the 9th International Conference on Dependability and Complex Systems DepCoS-RELCOMEX*, June 30-July 4, 2014, Brunów, Poland / Wojciech Zamojski [et al.] (eds.) Springer, 2014
- [13] J. Yang, et al. (2009) A Profile-based Approach to Just-in-Time Scalability for Cloud Applications. *CLOUD 2009 - 2009 IEEE International Conference on Cloud Computing*, September 21, 2009
- [14] Q. Zhang, et al. Cloud Computing: State-of-the-art and Research Challenges. *Journal of Internet Services and Applications*. 1, 1 (2010) pp. 7-18
- [15] B-H Li, et al: Cloud Manufacturing: a New Service-oriented Networked Manufacturing Model. *Computer Integrated Manufacturing Systems CIMS* 2010; 16(1) pp. 1-7
- [16] S. Li, L. Xu, X. Wang, J. Wang, Integration of Hybrid Wireless Networks in Cloud Services Oriented Enterprise Information Systems, *Enterprise Information Systems*, Vol. 6, No. 2, 2012, pp. 165-187
- [17] Q. Li, Z. Wang, W. Li, J. Li, C. Wang, R. Du, Applications Integration in a Hybrid Cloud Computing Environment: Modelling and Platform, *Enterprise Information Systems*, Vol. 7, No. 3, 2013, pp. 237-271
- [18] L. Ren, L. Zhang, F. Tao, X. Zhang, Y. Luo, Y. Zhang, A Methodology towards Virtualization-based High Performance Simulation Platform Supporting Multidisciplinary Design of Complex Products, *Enterprise Information Systems*, Vol. 6, No. 3, 2012, pp. 267-290
- [1] F Tao, L Zhang, YF Hu. *Resource Services Management in Manufacturing Grid System*. Wiley, Scrivener Publishing, 2012, ISBN 978-1-118-12231-0
- [2] F. Tao, L. Zhang, V. C. Venkatesh, Y. Luo, Y. Cheng: *Cloud Manufacturing- a Computing and Service-oriented Manufacturing Model*.

Proceedings of the Institution of Mechanical Engineers, Part B: Journal of Engineering Manufacture, 225(10), 2011, pp. 1969-1976

- [3] X. Xu: From Cloud Computing to Cloud Manufacturing, Robotics and Computer-Integrated Manufacturing, Volume 28, Issue 1, February 2012, pp. 75-86

Application of ANFIS for the Estimation of Queuing in a Postal Network Unit: A Case Study

Bojan Jovanović¹, Tatjana Grbić¹, Nebojša Bojović², Momčilo Kujačić¹, Dragana Šarac¹

¹ University of Novi Sad, Faculty of Technical Sciences

Trg Dositeja Obradovića 6, 21000 Novi Sad, Serbia

bojanjov@uns.ac.rs, tatjana@uns.ac.rs, kujacic@uns.ac.rs, dsarac@uns.ac.rs

² University of Belgrade, Faculty of Transport and Traffic Engineering

Vojvode Stepe 305, 11000 Belgrade, Serbia

nb.bojovic@sf.bg.ac.rs

Abstract: Regardless the level of technological development in a community, the unavoidable phenomenon is the appearance of queuing. This situation will continue as long as there is the customer's need for the direct contact with the service suppliers, as the case is in the Post of Serbia. The aim of the paper is to estimate the time a customer will spend in queuing while approaching the counter for financial services in a postal network unit. The observed system comprises a single queue, three handling channels and the service according to the FIFO principle. This paper presents a developed model that is realized in the following phases: recording data, preparing data for training, training the neuro-fuzzy system, forming a data set for testing where the expected mean service speed is obtained using the moving average method, and testing the neuro-fuzzy model. Observing the mass service system has so far been directed towards, the evaluation of their behaviour in the past which presents a basis to evaluate whether the system provides satisfactory performances. This paper moves a step in the direction of the behaviour evaluation of a mass service system, in the future, in order to observe whether it is possible to predict the service quality level to be provided to a customer. System customer in this case is not limited by the number of demanded services.

Keywords: Waiting time; Postal network unit; Financial services; ANFIS

1 Introduction

Inspired by the fact that the contemporary research of the mass service systems has been directed towards the system analysis in the past, as well as the significance of the system behaviour in the future and the possible application in the estimated queuing time (for providing the adequate level of service and the

engagement of the necessary resources), we have decided to apply the combination of all the existing methods to a postal network unit in the Post of Serbia. This paper presents what can be considered an innovative method presenting the combination of the two known methods, ANFIS approach and moving average method. Previous approaches were only based on one method, e.g. on the mass service theory or Monte Carlo method, see [1, 26]. Since the method presented in the paper introduces the combination of the ANFIS approach and the moving average method, we expect that this method or a modification of it will find the application both in other postal units and everywhere else where queuing is a regular phenomenon. In order to realize the mass service system management in the real time, the dataset that trained the model in ANFIS has been modified using the moving average method with the goal to estimate the system behaviour, which can be considered a step forward in comparison to previous research.

Mass service system managers are faced with the issue of providing an optimal relationship between the engaged resources and the time spent by the customer in order to realize their own demands. The satisfaction with the waiting time is not only a determinant for the satisfaction with the service; it also moderates the relation satisfaction – loyalty [6].

The traditional standpoint in the mass service theory is aimed at the optimal determination of the handling channel, where the customer is seen as an external factor. However, more attention is devoted recently to customer behaviour and unification of behavioural aspects, as well as studies dealing with service operations. The impact of waiting on the evaluation of service presents the main focus in numerous papers [5, 9, 18, 28, 31].

The question for queue management is not only the real time spent by a customer in a queue, it is also the customer's perception of that time and the related level of their satisfaction with the service [8]. Information on the waiting time has a positive impact on the customer's satisfaction with the waiting time [6].

After the realized service, the customer knows the waiting time spent and they use the experience to adjust their view on the average waiting time in a visited facility. Following it, one can present the customer's degree of satisfaction as a quotient of the acceptable and the expected waiting time based on the previous experience [30]. If a customer is provided with the information on the expected waiting time, their acceptable waiting time will be close to the expected time, since that information would influence the decision whether to approach the system or not at a certain period of time.

The problem that service providers have to deal with and that leads to alternate service quality is observed in the fluctuation of service demands. Solving the problem can lead, in one hand, to the development of more flexible systems (during the peak load additional resources are included), or on the hand, it can influence the customer's motivation by providing more satisfactory service

realization conditions during low demand periods [6]. Likewise, coordinating demands and capacities will allow for the utilization of diverse strategies for the queue organization, considering their configuration or the formation of a certain ambient, so that waiting can be more interesting and more bearable [34].

When waiting time estimation is considered, the research in call centres has been directed towards queuing in the FIFO system [3, 10, 11, 14, 32].

Services realized by traditional postal operators are characterized by heterogeneity, which is a consequence of different contents of postal items, as well as a wide range of financial services (payments out and payments in from bank accounts, agents work for the benefit of banks, advice on receipts, etc.). The demand for services provided by traditional postal operators has a stochastic character and it cannot be accurately determined since it fluctuates in intensity over time and space. These oscillations regarding service demands hamper technological processes in terms of providing adequate resources for their implementation.

According to the Law on payment transactions [25], the public postal operator in Serbia is allowed, in addition to providing postal services, to provide financial services as well. The banking sector as a competition to the public postal operator in the field of financial services is concentrated mainly on the segment of customers who are creditworthy. It is very important to include the segments of population that are not of interest to the banks. Considering this, the state has the venue for the financial inclusion through the public postal operator. As a consequence, the financial availability is reflected in adjusting prices to the solvency of the population segments that are the target area of the financial inclusion. The prediction of the waiting time in postal network units is of great importance in order to ensure adequate capacity for providing these services. Likewise, in order to offer a higher quality to the financial service beneficiaries, it is essential for them as well to be provided with the information on the queuing time, which is the research topic discussed further in the paper.

The paper is organized in the following manner: second section contains the analysis of the existing conditions in a postal network unit, the third section contains a proposition for a model for waiting time estimation, and the last section provides the conclusion with the directions for future research.

2 Condition Analysis in a Postal Network Unit

Data acquisition was conducted by recording queuing in a postal network unit for providing services to customers. The observed postal unit included counters for financial services, as well as counters for the reception and delivery of postal items. The selected postal unit is located in a residential neighbourhood, as well as

in the area with a wealthy number of diverse contents (green market, shopping mall and school). Accordingly, beneficiaries were both the people living in the area, as well as visitors to the green market, shopping mall, etc. Based on the above, it can be considered that the observed postal unit is a representative unit for providing customer service since it is not focused only on a specific segment of customers and since the users of financial services form a single queue.

The objective of the research refers to the counters for the implementation of financial services, so they were in the focus of the recording. Considering the fact that the volume of financial services is the highest in March and December [22, 23, 24], the recording was completed in the period of two weeks in December (from December, 6 to December, 18), i.e. the period of 12 working days. Working hours of the post office is from 7:00 am to 7:00 pm on weekdays and from 7:00 am to 2:00 pm on Saturdays. The observed system includes three handling channels, with a single queue and the service carried out according to the FIFO principle. The recording included the number of 5727 system customers. Following the Kendall's notation queuing theory, the system can be classified as M/G/3.

Testing the input stream of customers was realized according to days in Statistica 10 (StatSoft software package). An example for the first day is provided in Table 1. It can be observed that the widths of the classes are defined in the periods of 50 seconds. Table column *Observed Frequency* presents the recording frequencies, while the column *Expected Frequency* presents the theoretical frequencies of the observed classes for exponential distribution.

Table 1
Testing the compatibility of the input stream with the exponential distribution

Upper Boundary	Variable: Var1, Distribution: Exponential Chi-Square = 10.11028, df = 5 (adjusted), p = 0.07217								
	Observed Frequency	Cumulative Observed	Percent Observed	Cumul.% Observed	Expected Frequency	Cumulative Expected	Percent Expected	Cumul.% Expected	Observed-Expected
<= 50.00000	294	294	53.84615	53.8462	274.1121	274.1121	9.194812	50.2037	19.8879
100.00000	119	413	21.79487	75.6410	136.4977	410.6098	4.578679	75.2033	-17.4977
150.00000	65	478	11.90476	87.5458	67.9709	478.5806	2.280014	87.6521	-2.9709
200.00000	29	507	5.31136	92.8571	33.8470	512.4276	1.135363	93.8512	-4.8470
250.00000	15	522	2.74725	95.6044	16.8546	529.2822	0.565369	96.9381	-1.8546
300.00000	9	531	1.64835	97.2527	8.3930	537.6751	0.281533	98.4753	0.6070
350.00000	10	541	1.83150	99.0842	4.1794	541.8545	0.140193	99.2408	5.8206
400.00000	1	542	0.18315	99.2674	2.0812	543.9357	0.069811	99.6219	-1.0812
450.00000	2	544	0.36630	99.6337	1.0364	544.9721	0.034763	99.8117	0.9636
500.00000	2	546	0.36630	100.0000	0.5161	545.4881	0.017311	99.9062	1.4839
< Infinity	0	546	0.00000	100.0000	0.5119	546.0000	0.017170	100.0000	-0.5119

The header of the table presents the values of χ^2 test, the number of degrees-of-freedom, and the resulting p – value. Based on the provided p-value of 0.07217, it can be observed that the distribution of the input stream for the first day corresponds to the exponential distribution at the significance level of 0.05. Figure 1, presents graphical interpretations of recorded data for the presented day.

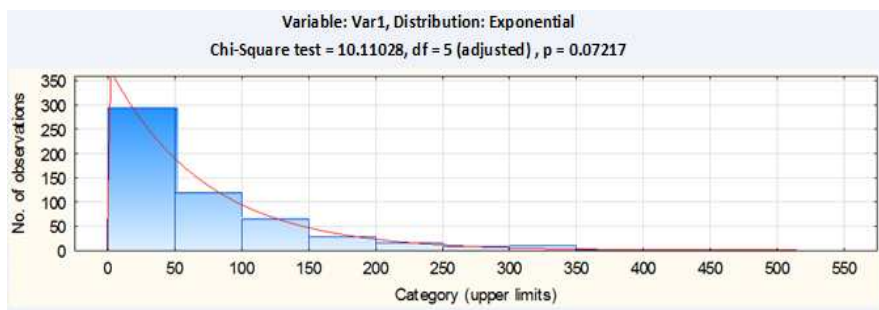


Figure 1

Graphical overview of the input stream of clients during the first day

Investigations were carried out for the remaining days as well, and the results obtained are shown in Table 2. Based on the results shown in Table 2, it can be concluded that the mid-intervals of the input streams of customers (time flow between successive arrivals of customers) during the observed days in nine cases correspond to the exponential distribution with the reliability level of 95%, i.e. p-values are lower than 0.05 in three cases out of the observed 12 days.

Table 2

p-values of the input streams of customers in testing the compatibility with exponential distribution

Days	1	2	3	4	5	6	7	8	9	10	11	12
p-values	0.072	0.428	<10 ⁻⁵ *	0.23	0.12	0.015*	0.174	0.161	0.55	0.662	0.186	2·10 ⁻⁴ *

By integrating the acquired data for the observed period, it is clear that the time intervals of mid-arrivals in the input stream of clients correspond to the exponential distribution with the parameter $\lambda = 1.026$ customer/min, with the reliability level of 95% (i.e. p-value equals 0.0964).

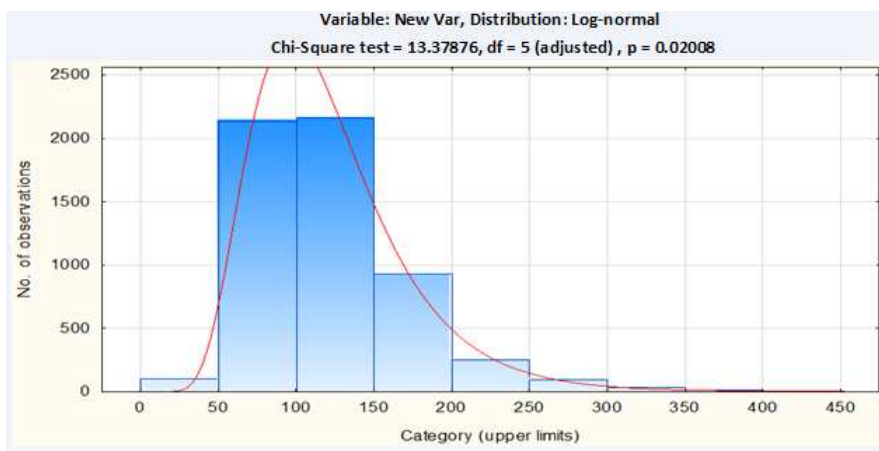


Figure 2

Distribution of service speed for the observed period of 12 days

After examining the compliance of the input data stream, the next to be tested was the service itself. It was established that the unified service for the observed period tends towards the log-normal distribution with the parameters $a = 4.722$ and $b = 0.153$ (Figure 2), with the reliability level of 99% since the low p-value of 0.02 was acquired. As for the data acquired for the observed days in 5 cases, after testing the compatibility with the log-normal distribution, it can be observed that the obtained p-values are lower for 0.01 (Table 3), whereas in other cases it can be considered that the service level corresponds to the log-normal distribution.

Table 3
p-values of the input streams of customers in testing the compatibility with log-normal distribution

Days	1	2	3	4	5	6	7	8	9	10	11	12
p-values	0.042	0.338	0.099	0.076	0.051	0.023	0.007*	$<10^{-5}$ *	$<10^{-4}$ *	0.001*	0.027	$<10^{-5}$ *

We can say that the precise analysis on the system with the form M/G/n is extremely demanding since the non-Markovian processes are happening inside the system, hence the formulas for calculating the system parameters that are widely applied in the cases of, for example, M/M/n systems, cannot be applied in this case. The observed system is characterized by heterogeneity, which is a consequence of different contents of postal items, as well as a wide range of financial services (payments out and payments in from bank accounts, agents work for the benefit of banks, advice on receipts, etc.). The demand for services provided by traditional postal operators has a stochastic character and it cannot be accurately determined since it fluctuates in intensity over time and space. In practice, the most commonly applied analysis is the approximation or simulation (as in the case of call centres [33, 35]). Concerning the observed system, the situation is even more complicated with the number of active handling channels altered during the observed interval. Differently from call centres, postal clerks are in a direct contact with service users and the overall ambient created in the counter hall. As a consequence, the working performances of clerks are exposed to the influence of the phenomenon of congestion which may lead to their working deterioration (clerks' fatigue) or improvement (striving to perform the job better to eliminate the occurrence of congestion). Likewise, the customer's possibility to claim an unlimited number of services further complicated the observance of the system. Considering the complexity of the observed system, i.e. the uncertainty and stochasticity as its properties, the research was focused on the development of a model based on the neuro-fuzzy approach for the evaluation of the waiting time.

3 Proposition for a Model

The main results are presented in this section of the paper. It describes the theoretical basis for the model created with ANFIS. After that, there is a proposition for a model for estimating the waiting time. The preparation of data for training ANFIS is realized. Training ANFIS is carried out by the fuzzy logic

toolbox, Matlab R2007b. It is followed by the data preparation for model testing. The test results are presented by RMSE (root mean square error), the coefficient of determination and the distribution of error values (differences between simulated and recorded values). At the end of this chapter, the average waiting time per each day is observed, as well as the movement of the queue length in selected days.

3.1 Neuro-Fuzzy Systems

Different techniques of artificial intelligence have been developed to solve problems of the real world using intelligent systems that possess skills similar to human skills in certain domains. Among them, fuzzy logic and neural networks are the most popular and widely applied in industrial applications [4, 15, 16, 17, 20, 21]. Fuzzy logic systems are widely applied in transportation engineering [29].

Adaptive Neuro-Fuzzy Inference System is a multi-layer adaptive network based on the fuzzy inference system [13]. ANFIS is a fuzzy inference system that can be trained on the basis of the acquired input-output data. The training method allows the system to adjust its parameters in order to perceive the input-output relationship hidden in the data set. Since it comprises two approaches (neural networks and fuzzy modelling), the appropriate inference in the quality and the quantity can be achieved [2].

In the case of the Sugeno fuzzy model of the first type with two inputs x and y , the rules are given in the following form:

Rule 1: if x is A_1 and y is B_1 then $f_1 = p_1x + q_1y + r_1$

Rule 2: if x is A_2 and y is B_2 then $f_2 = p_2x + q_2y + r_2$ (1)

where A_1, A_2, B_1, B_2 are membership functions for the inputs x and y respectively, whereas $p_1, q_1, r_1, p_2, q_2, r_2$, are the parameters of the output functions [27]. The corresponding ANFIS structure consists of five layers that are implemented by diverse functions of the nodes for training and by adjusting the parameters of a fuzzy system (Figure 3). Nodes located within the same layer have similar functions. The output of the i -th node within the layer l can be referred to as $O_{l,i}$. The functioning of the presented ANFIS system can be graphically represented as follows [13]:

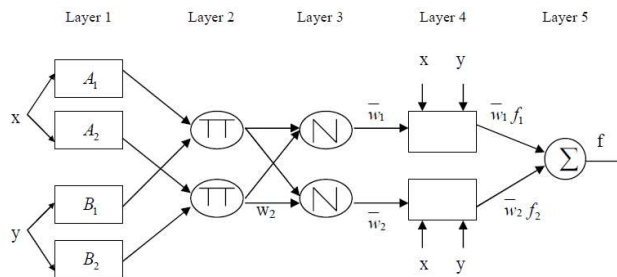


Figure 3

ANFIS architecture of the Sugeno fuzzy model with two inputs [13]

- Layer 1: The nodes of this layer generate membership functions for input variables. The output of the nodes $O_{1,i}$ is defined by the following expressions:

$$\begin{aligned} O_{1,i} &= \mu_{A_i}(x) & \text{for } i = 1, 2 \text{ or} \\ O_{1,i} &= \mu_{B_{i-2}}(y) & \text{for } i = 3, 4 \end{aligned} \quad (2)$$

where x and y are inputs into the node, while A_i and B_i are fuzzy sets related to the observed node and are defined by the shape of their membership function. Membership functions A_i and B_i can be presented by the generalized “bell” function:

$$\mu_i(x) = \frac{1}{1 + [(x - c_i)/a_i]^{2b_i}} \quad (3)$$

where a_i , b_i and c_i are parameters that change the shape of the membership function $\mu_i(x)$ from the minimum value 0 to the maximum value 1. The parameters of this layer correspond to the parameters of the premise (hypothesis) of the fuzzy model. Outputs of the first layer are values of the membership function of the premise.

- Layer 2: The nodes labelled with Π make up the second layer, which means that the input signals in the node are multiplied and the output of the node $O_{2,i}$ presents a strength of the i -th rule w_i which is calculated as follows:

$$O_{2,i} = w_i = \mu_{A_i}(x) \mu_{B_i}(y) \quad i = 1, 2 \quad (4)$$

- Layer 3: In this layer the nodes labelled with N calculate the ratio of the strength of the i -th rule and the sum of strengths of other rules, while the normalized power of the i -th rule is obtained as follows:

$$O_{3,i} = \bar{w}_i = \frac{w_i}{w_1 + w_2} \quad i = 1, 2 \quad (5)$$

- Layer 4: The nodes of the fourth layer calculate the contribution of the i -th rule to the output of the system with the following node function:

$$O_{4,i} = \bar{w}_i f_i = \bar{w}_i (p_i x + q_i y + r_i) \quad i = 1, 2 \quad (6)$$

where $\bar{w}_i$ is the output of the third layer, while a set of parameters (p_i, q_i, r_i) corresponds to the parameters of consequences.

- Layer 5: This layer consists of one node that is denoted by Σ and calculates the total output of the ANFIS as follows:

$$O_{5,i} = \sum_i \bar{w}_i f_i = \frac{\sum_i w_i f_i}{\sum_i w_i} \quad (7)$$

It may be noted that within the ANFIS architecture there are two adaptive layers, those being the first and the fourth layer. The parameters that are set within the first layer are connected to the input membership function (in the explained example those are parameters a_i , b_i and c_i), the so-called premise parameters. Within the fourth layer, parameters that are set relate to the first-order polynomial (p_i , q_i and r_i) and are referred to as consequence parameters [12, 13].

Neuro-fuzzy inference system is optimized by adapting the premise parameters and the consequence parameters in a manner as to minimize the defined objective function (most common, the difference between model outputs and actual outputs). The methods for improving the ANFIS parameters may include the gradient descent and Least Square Error (LSE) [13]. Chen (1999) compared the algorithms for training parameters in ANFIS membership functions [7]. The paper applied the hybrid learning algorithm that is a combination of the least square estimation and the back-propagation algorithms [13].

3.2 Model Development and Results Overview

Developing the model that is presented in this paper has been implemented through the following stages (Figure 4): recording the mass service systems, creating a data set for training ANFIS, training ANFIS, establishing data for testing, and testing and evaluating the model. During the recording process, the following variables were included: the number of active handling channels, queue length (i.e. number of customers) and the speed of the service level (customer/min). These values are used as inputs to ANFIS.

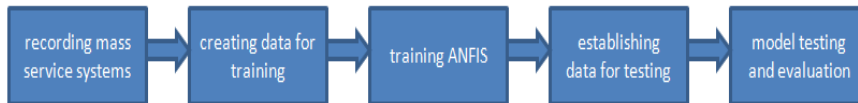


Figure 4

A model for the waiting time estimation

In preparing data for system training, it was necessary to modify the input size of the handling speed. Namely, the problem occurring is reflected in the fact that, when a customer accesses the queue, it is necessary to predict the service speed until the arrival to the service. Regarding this, the decision was to observe the service speed in segments of ten customers, for whom the average service speed was calculated. It was decided not to train the system with data on speed service that was recorded for each customer since in this manner the system would become too “sensitive” to this parameter. In other words, the prediction methods could not provide exact figures of the level obtained at the level of recording. The next phase in the observed model is the system training. The training of the system is realized in the ANFIS Editor GUI, within the software Matlab R2007b.

Applying the subtractive clustering method, with the default parameter values being the range of influence 0.5, squash factor 1.25, accept ratio 0.5 and reject ratio 0.15, the initial fuzzy training system was formed, i.e. Sugeno fuzzy model with three inputs (Figure 8). Each of the three inputs is associated with two membership functions (Figure 5, 6, 7). Membership functions in Figure 5, are defined by the following parameters: blue - $e^{-\frac{(x-1.968)^2}{0.25}}$ and red - $e^{-\frac{(x-2.032)^2}{0.2429}}$. Membership functions in Figure 6, are defined by the following parameters: blue - $e^{-\frac{(x-0.47)^2}{0.0561}}$ and red - $e^{-\frac{(x-0.6041)^2}{0.05}}$. And the parameters for the membership function in Figure 7, are as follows: blue - $e^{-\frac{(x-6.001)^2}{76.6322}}$ and red - $e^{-\frac{(x-17)^2}{76.5579}}$.

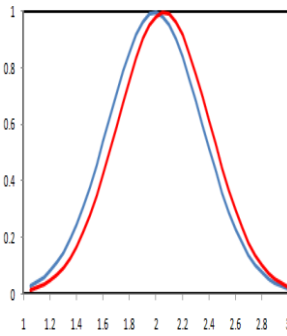


Figure 5
Number of active handling
channels

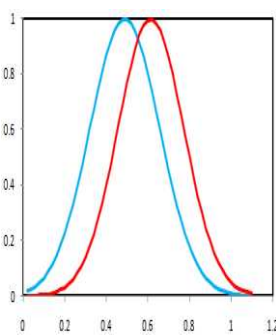


Figure 6
Service speed
(customer/min)

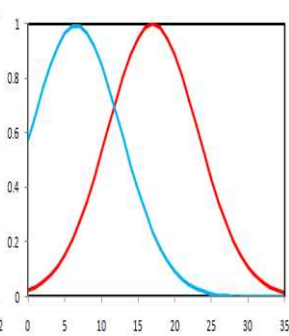


Figure 7
Queue length
(number of customers)

The system is defined by the two-rule base, as well as by two output membership functions. The number of epochs for training the system is 300, while the hybrid algorithm for training is implemented. Features of the system are stabilized after 20 epochs already to the RMSE value of 80 seconds.

Altering the values of the initial set parameters for applying the subtractive clustering method does not bring any significant improvement related to the RMSE, even with the increase in the number of the membership functions (Table 4). In the case of decreasing the value of the *range of influence*, the number of generated membership functions is increased. A little decrease in RMSE to 77 seconds is observed, though the surface of the transfer function is dramatically disrupted. The similar situation applies to the alteration of other parameters in the direction of increasing the number of the membership functions.

The graphical representation of the transfer function of the obtained system can be observed according to the pairs of input sizes: number of handling channels vs. service speed (Figure 9), number of handling channels vs. queue length (Figure 10), and service speed vs. queue length (Figure 11). A good feature of the obtained surfaces is reflected in the fact that they do not have any peaks, that is, there is a bland transition on changing the input values.

Table 4
RMSE as parameters in the subtractive clustering method

range of influence	squash factor	accept ratio	reject ratio	number mf	RMSE
0.5	1.25	0.5	0.15	2+2+2	80.79
0.4	1.25	0.5	0.15	3+3+3	81.07
0.3	1.25	0.5	0.15	5+5+5	78.91
0.2	1.25	0.5	0.15	10+10+10	78.14
0.1	1.25	0.5	0.15	33+33+33	77.16

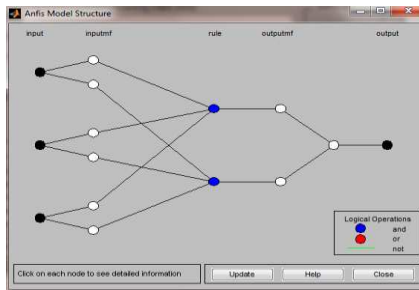


Figure 8
ANFIS structure of the proposed model

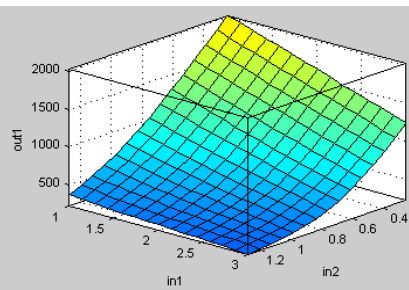


Figure 9
Transition function for the input of handling channels vs. service speed

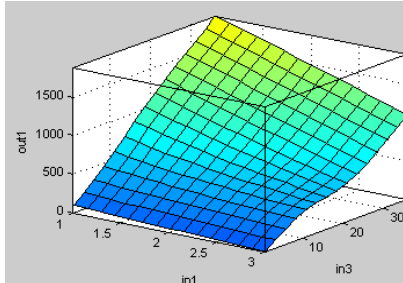


Figure 10
Transition function for the input of handling channels vs. queue length

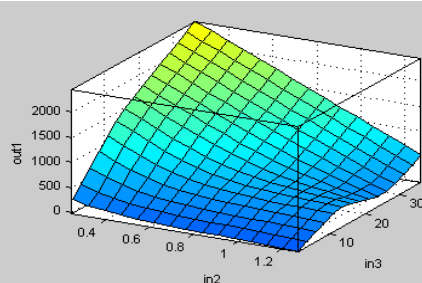


Figure 11
Transition function for the input of service speed vs. queue length

After training the model, a new data set for the system testing was established. In addition to input values, the number of handling channels and the queue length, the value to be estimated was the speed of the handling system when the n -th customer approaches. The estimation of the speed rate was achieved using the simple moving average method, in a manner that the speed from 3 previous intervals of ten recorded data, each was considered. Data set for testing was reduced to 5376 due to the initial segments for the moving average method for each day.

In order to examine the degree of similarity of the developed model with the recorded conditions in the postal network unit, absolute errors simulated vs recorded were calculated. The distribution of errors is given in Figure 12, where it can be observed that the expected value of the error is 0.451 seconds with a standard deviation of 154 seconds. The level of similarity of the simulated and the recorded data is expressed with the coefficient of determination (Figure 13):

$$R^2 = \left[\frac{\sum (S_r - \bar{S}_r)(S_s - \bar{S}_s)}{\sqrt{\sum (S_r - \bar{S}_r)^2 \sum (S_s - \bar{S}_s)^2}} \right]^2 \quad (8)$$

where S_r are the recorded values, while S_s are simulated values.

In comparison to the data recorded, the simulated data of the represented model have RMSE of 154 seconds and the coefficient of determination with the value of 0.826.

Likewise, the model simulation was implemented with the moving averages of 4 and 10 intervals, where the results obtained for RMSE were 158 sec and 167 sec, while the coefficients of determination were 0.817 and 0.796, respectively.

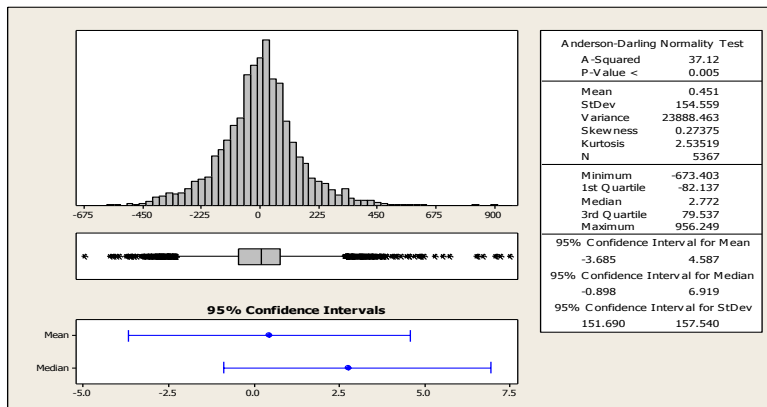


Figure 12
Distribution of the error values in Minitab 16

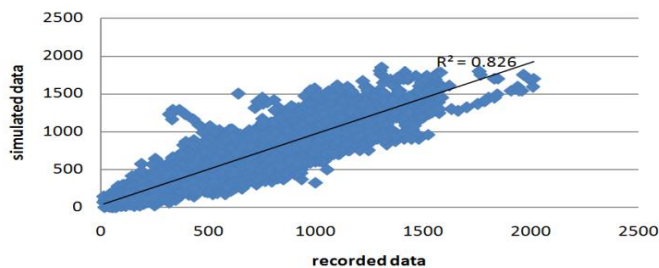


Figure 13
Comparison of recorded values and values obtained in ANFIS

The model used in the paper is the one with the moving averages from 3 previous intervals. It also possesses the advantage observed in the fact that too much time would not elapse before the model begins the estimation in comparison to the beginning of the working hours in the postal network unit for providing services to customers. Certainly, the sooner the estimation begins, a lower number of unsatisfied customers is to be expected since queuing appears to last longer to a customer with the lack of information on its duration compared to the customer with the known waiting time [19].

In continuation, the average waiting time per day was observed, as well as the values of average waiting times obtained by simulation. Table 5, provides the ratio of average waiting times that were recorded and average times obtained by simulation, where it can be observed that these values are very similar (maximum deviation is 17 seconds).

The minimum average waiting time of 176 seconds was observed on the fourth day of recording, i.e. in the ninth day of the month when different inquiries influencing the increase in volume of financial services have not yet come to their due dates. During this day, the peaks related to the queue length reached the value of 7 customers at 9:50 and 11:30 in the morning. In the afternoon, the queue length reached a maximum value of 3 customers at 14:00, 15:00 and 18:00.

Table 5
Average waiting times recorded vs. simulated

Day	1	2	3	4	5	6	7	8	9	10	11	12
Average waiting time, recorded, sec.	367	290	273	176	388	366	680	485	409	645	707	665
Average waiting type of the model, sec.	366	293	256	183	389	372	678	490	422	647	700	654

The highest average waiting time was observed on the 11th day of recording. The reason for this is reflected in the fact that it was the 17th day of the month when diverse inquiries are due to be paid. Additional impact was caused by the situation that it was the last workday for most customers, and also by the fact that business policies of some companies provide their customers a certain discount for payments before the 20th of the month. The maximum queue length in the morning reached a number of 24 customers at 10:30. In the afternoon, the longest queue was observed at 16:20, with 28 customers. As an average day, there is the 8th day of recording. In the morning, the queue length reached the maximum value of 27 customers at 11:30, while in the afternoon the maximum queue length was with 14 users at 13:45.

Conclusion

This paper presents a model for estimating the waiting time based on ANFIS. The developed model shows a satisfactory result considering the value of the coefficient of determination $R^2 = 0.826$, as well as the value of RMSE of 154

seconds. The average waiting time recorded per day and the average waiting time obtained by the proposed model provide similar results (Table 5). Estimating the service time as the input value required an adequate modification (average value of 10 serving times, as well as the application of the moving average method), so that the proposed model would be more reliable to describe the observed queuing system.

The contribution of the model is in providing information to both the customers and the management of the postal network unit. Providing information on the waiting time presents an additional quality of the service since customers can expect how much of their time to set aside, which makes the waiting process appear acceptable. The same information provides an opportunity for post office managers to manage a number of active counters in real time. In other words, by changing the parameters in the proposed system an adequate simulation can be implemented, which will then indicate whether a particular procedure is justified, i.e. whether the available resources are utilized in an optimal manner.

Further development will be focused on estimating the number of consumers who will ask for a service at specified time intervals and, thus, provide a timely activity in order to try to prevent the congestion of the system. Further considerations may include the number of services to be observed in order to obtain an average speed of service in the observed time interval, as well as the number of sample intervals in estimating the service speed by using the moving average method.

Acknowledgements

This research is supported by the Ministry of Science of Serbia, Grant Numbers 36040, 174009, 32035.

References

- [1] I. Adan, J. Resing, Queueing theory, Dept. of Mathematics and Computing Science, Eindhoven University of Technology, Netherlands, 2002
- [2] M. Alizadeh, F. Jolai, M. Aminnayeri, R. Rada, Comparison of Different Input Selection Algorithms in Neuro-Fuzzy Modeling, Expert Systems with Applications, 39(1), 1536-1544, 2012
- [3] M. Armony, C. Maglaras, Contact Centers with a Call-Back Option and Real-Time Delay Information, Operations Research, 52(4), 527-545, 2004
- [4] G. Athanasopoulos, C. R. Riba, C. Athanasopoulou, A Decision Support System for Coating Selection Based on Fuzzy Logic and Multi-Criteria Decision Making, Expert Systems with Applications, 36(8), 10848-10853, 2009
- [5] J. Baker, M. Cameron, The Effects of the Service Environment on Affect and Consumer Perception of Waiting Time: an Integrative Review and Research Propositions, Journal of the Academy of Marketing Science, 24(4), 338-349, 1996

- [6] F. Bielen, N. Demoulin, Waiting Time Influence on the Satisfaction-Loyalty Relationship in Service, *Managing Service Quality*, 17(2), 174-193, 2007
- [7] M. S. Chen, A Comparative Study of Learning Methods in Tuning Parameters of Fuzzy Membership Functions, In *Systems, Man, and Cybernetics, 1999 IEEE SMC'99 Conference Proceedings. 1999 IEEE International Conference on* (Vol. 3, pp. 40-44) IEEE, 1999
- [8] M. M. Davis, J. Heineke, Understanding the Roles of the Customer and the Operation for Better Queue Management, *International Journal of Operations & Production Management*, 14(5), 21-34, 1994
- [9] M. K. Hui, D. K. Tse, What to Tell Consumers in Waits of Different Lengths: An Integrative Model of Service Evaluation, *Journal of Marketing*, 60, 81-90
- [10] R. Ibrahim, W. Whitt, Real-Time Delay Estimation Based on Delay History, *Manufacturing & Service Operations Management*, 11(3), 397-415, 2009
- [11] R. Ibrahim, W. Whitt, Wait-Time Predictors for Customer Service Systems with Time-Varying Demand and Capacity, *Operations Research*, 59(5), 1106-1118, 2011
- [12] J. S. Jang, Self-Learning Fuzzy Controllers Based on Temporal Back Propagation, *Neural Networks, IEEE Transactions on*, 3(5), 714-723, 1992
- [13] J. S. Jang, ANFIS: Adaptive-Network-based Fuzzy Inference System, *Systems, Man and Cybernetics, IEEE Transactions on*, 23(3) 665-685, 1993
- [14] O. Jouini, Z. Aksin, Y. Dallery, Call Centers with Delay Information: Models and Insights. *Manufacturing & Service Operations Management*, 13(4), 534-548, 2011
- [15] C. Kahraman, S. Çevik, N. Y. Ates, M. Gülbay, Fuzzy Multi-Criteria Evaluation of Industrial Robotic Systems. *Computers & Industrial Engineering*, 52(4), 414-433, 2007
- [16] J. Liebowitz, *The Dynamics of Decision Support System and Expert System*, Orlando: The Dryden Press, 1990
- [17] J. Liebowitz, D. A. DeSalvo, *Expert Systems for Business and Management*, Yourdon Press, 1989
- [18] W. J. Luo, M. L. Liberatore, R. B. Nydick, Q. Chung, E. Sloane, Impact of Process Change on Customer Perception of Waiting Time: a Field Study, *Omega*, 32(1), 77-83, 2004
- [19] D. H. Maister, *The Psychology of Waiting Lines* (pp. 71-78) Harvard Business School, 1984
- [20] J. Mendes, R. Araújo, F. Souza, Adaptive Fuzzy Identification and Predictive Control for Industrial Processes, *Expert Systems with Applications*, 40(17), 6964-6975, 2013

- [21] M. Paliwal, U. A. Kumar, Neural Networks and Statistical Techniques: A Review of Applications, Expert Systems with Applications, 36(1), 2-17, 2009
- [22] Post Serbia, Company Profile 2009, available at: <http://www.posta.rs/dokumenta/eng/o-nama/profil-preduzeca/Company%20Profile%202009.pdf> 2010
- [23] Post Serbia, PTT glasnik, available at: <http://www.posta.rs/dokumenta/PTTglasnik/Postanski%20glasnik%20specijal%20%20april%202011.pdf>, 2011
- [24] Post Serbia, Company Profile 2011, available at: <http://www.posta.rs/dokumenta/eng/o-nama/profil-preduzeca/Company%20Profile%202011.pdf>, 2012
- [25] Republic Serbia, Law on Payment Transactions, FRY Official Gazette, Nos 3/2002 and 5/2003, and RS Official Gazette, Nos 43/2004, 62/2006 and 31/2011
- [26] R. Y. Rubinstein, D. P. Kroese, Simulation and the Monte Carlo Method (Vol. 707) John Wiley & Sons, 2011
- [27] T. Takagi, M. Sugeno, Fuzzy Identification of Systems and its Applications to Modeling and Control, Systems, Man and Cybernetics, IEEE Transactions on, (1) 116-132, 1985
- [28] S. Taylor, Waiting for Service: the Relationship between Delays and Evaluations of Service, The Journal of Marketing, 56-69, 1994
- [29] D. Teodorović, Fuzzy Logic Systems for Transportation Engineering: the State of the Art. Transportation Research Part A: Policy and Practice, 33(5), 337-364, 1999
- [30] A. van Ackere, C. Haxholdt, E. R. Larsen, Dynamic Capacity Adjustments with Reactive Customers. Omega, 41(4), 689-705, 2013
- [31] A. Whiting, N. Donthu, Managing Voice-to-Voice Encounters Reducing the Agony of Being Put on Hold, Journal of Service Research, 8(3), 234-244, 2006
- [32] W. Whitt, Predicting Queuing Delays, Management Science, 45(6), 870-888, 1999
- [33] W. Whitt, Engineering Solution of a Basic Call-Center Model. Management Science, 51(2), 221-235, 2005
- [34] V. Zeithaml, J. Bitner, Services Marketing: Integrating Customer Focus across the Firm, 3rd ed., McGraw-Hill Higher Education, New York, 2002
- [35] S. Zeltyn, A. Mandelbaum, Call Centers with Impatient Customers: Many-Server Asymptotics of the M/M/n+ G Queue. Queueing Systems, 51(3-4), 361-402, 2005

Application of Fuzzy and Possibilistic *c*-Means Clustering Models in Blind Speaker Clustering

Gábor Gosztolya¹, László Szilágyi^{2,3}

¹ MTA-SZTE Research Group on Artificial Intelligence of the Hungarian Academy of Sciences and University of Szeged, Tisza Lajos krt. 103, H-6720 Szeged, Hungary, ggabor@inf.u-szeged.hu

² Dept. of Control Engineering and Information Technology, Budapest University of Technology and Economics, Hungary

³ Computational Intelligence Research Group, Dept. of Electrical Engineering, Sapientia University of Transylvania, Tîrgu Mureş, Romania, lalo@ms.sapientia.ro

*Abstract: Blind Speaker Clustering is a task within speech technology, where we have a collection of speech recordings (utterances), and the goal is to identify which utterances belong to the same speakers. To aid the clustering process in this task, we performed pre-processing steps such as feature selection and Principal Component Analysis (PCA); still, the choice of clustering method is not a trivial one. To find the best performing algorithm, we tested standard methods such as *k*-means (or hard *c*-means, HCM) and fuzzy *c*-means (FCM) as well as several improved versions of FCM. In the end, we were able to achieve the best performance using probabilistic-possibilistic mixture partitions. The obtained purity score of 83.9% is significantly higher than the baseline score of 46.9%.*

*Keywords: clustering; fuzzy *c*-means algorithm; possibilistic *c*-means algorithm; speech technology*

1 Introduction

Automatic Speech Recognition (ASR) seeks to create the correct transcription (written form) of an utterance (a recording containing speech). Traditionally, speech technology researchers focused primarily on ASR (e.g. [1-3]), but in the last few years another area has received growing attention. It is called computational paralinguistics, and it seeks to extract non-verbal information from the speech signal. This area includes tasks such as emotion detection [4, 5], speaker age estimation [6], conflict intensity estimation [7-9], detecting social signals like laughter and filler events [5], and estimating the amount of physical or

cognitive load during speaking [10-12]. Several of these tasks attempt to detect phenomena which vary from speaker to speaker. Therefore, in these tasks, if we could identify which utterances (recordings) belong to the same person, it would clearly assist the following classification or regression step. This task, called Blind Speaker Clustering (or just Speaker Clustering), is a viable tool in computational paralinguistics; for example, it was shown that speaker-wise data normalization can lead to a significantly improved classification performance compared to using global normalization techniques [10, 13].

Speaker clustering is an existing, current problem in speech recognition literature (e.g. [14-16]). In most cases, however, it has to be performed along with speaker segmentation ("Who spoke when?"), while now we have only one speaker in an utterance; hence we have to concentrate only on speaker clustering. Note that an important aspect of this task is that we have to work with the utterances of speakers that are unseen to us at training time, so it is clearly a clustering task.

Clustering means that we form groups of those examples which are similar to each other and different from the others, but this definition does not tell us *in which sense* they are different. For example, speech utterances may be similar if they record the speech of the same speaker, or the speakers utter the same sentences, or they were recorded under similar conditions (microphone, background noise), etc. In the actual task, however, we would like to separate the different speakers. A straightforward choice to control the way of clustering is by applying feature selection. If we keep only those attributes which correlate well with the desired property of examples, we can control the type of clusters formed. Still, we have to keep in mind the fact that we also have to avoid choosing a redundant feature subset, as it can also hinder the clustering process.

Besides feature selection, an important choice is that of the clustering method. A straightforward choice is the k -means (or hard c -means, HCM, [17]) algorithm; however, it is a stochastic algorithm, which is vulnerable to random initialization. To this end, we also test fuzzy c -means (FCM [18]) in this task, as well as three of its improved variants ([9, 19, 20]).

2 Blind Speaker Clustering by Feature Selection

Blind Speaker Clustering can be simply viewed as a clustering problem, for which standard clustering methods such as the HCM and FCM algorithms can be readily applied. However, these methods have a weak point, namely that they work in a multi-dimensional space treating all dimensions as equally important (as they rely on the Euclidean distance of the points). This means that they are sensitive to differently-scaled, redundant and irrelevant features.

The first issue can be handled by normalization, i.e. all the features can be normalized (i.e. scaled to a fixed interval, e.g. $[0, 1]$ or $[-1, 1]$) or standardized (i.e. transformed so as to have zero mean and unit standard deviation). The other two issues can be handled via feature selection; in fact, we can turn the feature selection to our advantage so that we just keep those features which help us create the right kind of clusters (in our case, different speakers).

To be able to perform this, we will need several things. First, to measure which feature set allows us to form better clusters, we will need a set of recordings with their correct classes (now: speakers) annotated. Second, we will have to choose an evaluation metric by which we will rank the results of the different clustering outcomes. Third, we will need to choose (or construct) a feature selection method.

2.1 Performing Feature Selection

A wide range of feature selection algorithms exist (e.g. [21, 22]); most of them, however, have a high computational complexity. To this end, we applied some quite quick and simple pre-processing steps to perform feature selection.

Feature selection has to deal with two phenomena, namely irrelevant and redundant features. In the first case, the problem is that some features do not assist the forming of the desired clusters, or even distract the clustering algorithms (e.g. describe relations that the actual speaker mentioned, and not *who* spoke in that given utterance). Yet, the redundant features describe the same phenomenon in a very similar way. As most clustering algorithms treat each feature as an equally important dimension, redundant features will have a larger importance overall, hence will distract the clustering method used.

2.1.1 Handling Irrelevant Features

We handled the issue of irrelevant features by applying a simple feature selection method. We took the feature vectors of two speakers, and calculated the correlation between each feature with the change of speakers. We repeated this for each speaker pair, and the absolute values of the resulting correlation values were averaged out. Then, the features were sorted according to their averaged correlation score in descending order, and we selected the most correlated features. This way, we also had control over the type of clusters formed.

2.1.2 Handling Redundant Features

The issue of redundancy was dealt with by using Principal Component Analysis (PCA, [23]). PCA is a statistical method which transforms our observation vectors into a space described by linearly uncorrelated directions (the principal components) via an orthogonal transformation. That is, the first direction returned

by the PCA will point to the direction where the variance of our data is the highest; the further directions will point to the directions which have the largest possible variance, provided that they are orthogonal to all previous directions. By transforming our examples into this coordinate system, and then performing normalization, we can get rid of most of the redundancy in our attributes.

PCA also supplies information about the importance of each new dimension (feature) describing the examples. It is common practice to keep only the first directions which describe at least a given amount (e.g. 90%) of the information stored in the example set, thereby applying PCA as a feature extraction tool [24]. Of course, the amount of information to be retained is not a trivial one, hence we experimented with different thresholds for this value as well.

3 The Employed Clustering Methods

Given a set of feature vectors $\mathbf{X} = \{ \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \}$ describing n objects, the fuzzy c -means (FCM) algorithm can produce a fuzzy partition into a predefined number of clusters c , based on the minimization of the quadratic objective function

$$J_{FCM} = \sum_{i=1}^c \sum_{k=1}^n u_{ik}^m \|\mathbf{x}_k - \mathbf{v}_i\|^2, \quad (1)$$

under the probabilistic constraint $\sum_{i=1}^c u_{ik} = 1, \forall k = 1, \dots, n$, where $\mathbf{v}_i$ represents the prototype or centroid of cluster i ($i = 1, \dots, c$), $u_{ik} \in [0, 1]$ is the fuzzy membership function showing the degree to which the vector $\mathbf{x}_k$ belongs to cluster i , and $m > 1$ is the fuzzyfication parameter. The minimization of the objective function J_{FCM} is achieved by alternately applying the optimization of J_{FCM} over u_{ik} with $\mathbf{v}_i$ fixed, $i = 1, \dots, c$, and the optimization of J_{FCM} over $\mathbf{v}_i$ with u_{ik} fixed, $i = 1, \dots, c, k = 1, \dots, n$ [18]. In each loop, the optimal values are deduced from the zero gradient conditions using Lagrange multipliers, and are computed using the following formulas:

$$u_{ik}^* = \frac{\|\mathbf{x}_k - \mathbf{v}_i\|^{-2/(m-1)}}{\sum_{j=1}^c \|\mathbf{x}_k - \mathbf{v}_j\|^{-2/(m-1)}} \quad \begin{array}{l} \forall i = 1, \dots, c \\ \forall k = 1, \dots, n \end{array} \quad (2)$$

$$\mathbf{v}_i^* = \frac{\sum_{k=1}^n u_{ik}^m \mathbf{x}_k}{\sum_{k=1}^n u_{ik}^m} \quad \forall i = 1, \dots, c. \quad (3)$$

According to the alternating optimization (AO) scheme of the FCM algorithm, eqs. (2) and (3) are alternately applied, until the cluster prototypes stabilize. This stopping criterion compares the sum of norms of the variations of the prototype vectors $\mathbf{v}_i$ within the latest iteration, with a predefined small threshold value ε . The algorithm requires proper initialization of the cluster prototypes. In our case, we have assigned randomly chosen input vectors to cluster prototypes, and ensuring that $\forall i, j \in \{1, \dots, c\}$ and $i \neq j$ we have $\mathbf{v}_i \neq \mathbf{v}_j$.

Hard c -means [17] is a special case of FCM, when $m \rightarrow 1$ and thus the memberships are obtained according to the winner-takes-all rule. Each cluster prototype is obtained as the mean of input vectors assigned to the given cluster. The first time when the partition does not change during an iteration, the convergence is achieved.

3.1 FCM Variants Employed in this Study

Several families and variants of c -means clustering models have been introduced recently, which reportedly produce better partitions than FCM in several applications. In the following, we enumerate those applied in our study.

3.1.1 Adding Possibilistic Component to Fuzzy c -means Clustering

The **possibilistic c -means clustering (PCM)** algorithm assigns typicality values to fuzzy membership functions [25]. Thus in PCM, the elements of the partition matrix, denoted by t_{ik} instead of u_{ik} ($i = 1, \dots, c$, $k = 1, \dots, n$), describe how compatible the input vectors are with the clusters represented by the computed cluster prototypes. Typicality values with respect to one cluster do not depend on any of the prototypes of other clusters.

Since PCM often produces coincident clusters, Pal et al. introduced a mixture clustering model called possibilistic-fuzzy c -means (PFCM) clustering that comprises a probabilistic and a possibilistic term [19]. PFCM optimizes the objective function:

$$J_{PFCM} = \sum_{i=1}^c \sum_{k=1}^n [a u_{ik}^m + b t_{ik}^p] \|\mathbf{x}_k - \mathbf{v}_i\|^2 + \sum_{i=1}^c \eta_i \sum_{k=1}^n (1 - t_{ik})^p, \quad (4)$$

where the fuzzy membership functions u_{ik} ($i = 1, \dots, c$, $k = 1, \dots, n$) are constrained by the probabilistic conditions, while the typicality values $t_{ik} \in [0, 1]$ ($i = 1, \dots, c$, $k = 1, \dots, n$) are subject to: $0 < \sum_{i=1}^c t_{ik} < c$, $\forall i = 1, \dots, c$. The fuzzy exponent m and possibilistic exponent p must be greater than 1, while a and b are tradeoff parameters to set the balance between the probabilistic and possibilistic term. The variables η_i ($i = 1, \dots, c$) are called possibilistic penalty terms and control the

variance of the clusters. The optimization formulas applied in each loop of the alternating optimization are:

$$t_{ik}^* = \left[1 + \left(\frac{b \|\mathbf{x}_k - \mathbf{v}_i\|^2}{\eta_i} \right)^{1/(p-1)} \right]^{-1} \quad \begin{array}{l} \forall i = 1, \dots, c \\ \forall k = 1, \dots, n \end{array} \quad (5)$$

$$\mathbf{v}_i^* = \frac{\sum_{k=1}^n [a u_{ik}^m + b t_{ik}^p] \mathbf{x}_k}{\sum_{k=1}^n [a u_{ik}^m + b t_{ik}^p]} \quad \forall i = 1, \dots, c. \quad (6)$$

The probabilistic part of the partition is computed exactly the same way as in FCM, according to Eq. (2). This algorithm was found to be robust in several tests.

3.1.2 Fuzzy-Possibilistic Product Partition c-means

The **fuzzy-possibilistic product partition c-means (FPPPCM)** algorithm was introduced with the goal to eliminate the outlier sensitivity of previous mixture clustering models [20]. This partition also employs a probabilistic and a possibilistic term, but it combines them via multiplication instead of via linear combination. The algorithm optimizes the objective function:

$$J_{FPPPCM} = \sum_{i=1}^c \sum_{k=1}^n u_{ik}^m \left[t_{ik}^p \|\mathbf{x}_k - \mathbf{v}_i\|^2 + (1 - t_{ik})^p \eta_i \right], \quad (7)$$

constrained by the conventional probabilistic and possibilistic conditions mentioned above. The only parameters of FPPPCM are the fuzzy exponent $m > 1$, the possibilistic exponent $p > 1$, and the conventional penalty terms of the possibilistic partition denoted by η_i , $i = 1, \dots, c$. The optimization formulas that stem from zero gradient conditions using Lagrange multipliers are:

$$t_{ik}^* = \left[1 + \left(\frac{\|\mathbf{x}_k - \mathbf{v}_i\|^2}{\eta_i} \right)^{1/(p-1)} \right]^{-1} \quad \begin{array}{l} \forall i = 1, \dots, c \\ \forall k = 1, \dots, n \end{array} \quad (8)$$

$$u_{ik}^* = \frac{\left[t_{ik}^p \|\mathbf{x}_k - \mathbf{v}_i\|^2 + \eta_i (1 - t_{ik})^p \right]^{-1/(m-1)}}{\sum_{j=1}^c \left[t_{jk}^p \|\mathbf{x}_k - \mathbf{v}_j\|^2 + \eta_j (1 - t_{jk})^p \right]^{-1/(m-1)}} \quad \begin{array}{l} \forall i = 1, \dots, c \\ \forall k = 1, \dots, n \end{array} \quad (9)$$

$$\mathbf{v}_i^* = \frac{\sum_{k=1}^n u_{ik}^{m_t p} \mathbf{x}_k}{\sum_{k=1}^n u_{ik}^{m_t p}} \quad \forall i = 1, \dots, c. \quad (10)$$

This algorithm has its main advantage of having a reduced number of parameters. It was found to efficiently reject the effect of outliers while being accurate also in the absence of outliers.

3.1.3 Suppressed FCM

Suppressed FCM was introduced with the intent of combining the quick convergence of HCM with the fine partitions produced by FCM. It manipulates with fuzzy membership functions produced by the FCM algorithm using Eq. (2): for each vector $\mathbf{x}_k$ it looks for the closest cluster prototype, say $\mathbf{v}_w$, applies suppression by multiplying all u_{ik} values by a suppression rate $\alpha \in [0,1]$, and increases u_{wk} by $1 - \alpha$ to maintain the probabilistic constraint [26]. These modified fuzzy membership values are then fed to Eq. (3) to update the cluster prototypes. The algorithms obtained for various values of α are reportedly quick and accurate in most clustering problems [27].

4 Experimental Setup

Next we will describe the way our experiments were performed: the way clustering accuracy was measured, the database used, the feature set extracted from the examples, and the way the parameters of the clustering methods were set.

4.1 Evaluation Metrics

If the real groups of examples (in our case, the different speakers) are known, we can evaluate a clustering hypothesis generated via an automatic clustering method (*external evaluation*, [28]). However, this is more difficult to do than for classification, as we cannot be sure which resulting cluster corresponds to which group (if any). Perhaps this is why there are several evaluation metrics available for this purpose.

One of the metrics that can be used for clustering evaluation is purity; this metric takes the most frequent class label in each cluster, and calculates the ratio of the elements in the cluster which belong to this class [28-30]. Then, these scores are averaged out for all clusters by weighting them with the number of their elements.

That is, for $\Omega = \{\omega_1, \dots, \omega_c\}$ (the set of resulting clusters), $C = \{\xi_1, \dots, \xi_N\}$ (the set of real groupings) and n elements ($\sum |\omega_j| = \sum |\xi_i| = n$), we calculate

$$Purity(\Omega, C) = \frac{1}{n} \sum_{j=1}^c \max_i |\omega_j \cap \xi_i|. \quad (11)$$

Bad clustering has a purity value close to zero, while a perfect clustering has a purity score of one. It has the drawback that it is easy to achieve high purity scores when the number of clusters (c) is large, but as in our case this is known in advance, we can set $c = N$ (the number of speakers) and handle this problem.

Another possibility is to use entropy [28, 30], which is defined as

$$E(\omega_j) = -\frac{1}{\log N} \sum_{i=1}^N \frac{|\omega_j \cap \xi_i|}{|\xi_j|} \log \frac{|\omega_j \cap \xi_i|}{|\xi_j|} \quad (12)$$

for any $j = 1, \dots, c$, and the entropy of the C clustering will be the sum of the $E(\omega_j)$ values weighted by the number of the elements. That is,

$$Entropy(\Omega, C) = \sum_{j=1}^c \frac{|\omega_j|}{n} E(\omega_j). \quad (13)$$

The better a clustering, the lower the entropy value it has; a perfect clustering has zero entropy.

4.2 The Munich Biovoice Corpus

We performed our experiments on the Munich Biovoice Corpus (MBC, [31]). It contains the utterances of 19 subjects (4 female and 15 male) of three nations (Chinese, German and Italian) both after light and heavy physical load. They had to pronounce sustained vowels as well as reading a short story, which was recorded by two different microphones. Besides the audio recordings, heart rate and skin conductivity was monitored as well. The dataset was later used in the Interspeech ComParE 2014 Physical Load Sub-Challenge [10].

4.3 Experimental Setup

In our experiments we employed the feature set used in [10]. It contained 6373 features overall, extracted by using the tool called openSMILE [32]. The set includes energy, spectral, cepstral (MFCC) and voicing related low-level descriptors (LLDs), as well as a few other LLDs including logarithmic harmonic-to-noise ratio (HNR), spectral harmonicity, and psychoacoustic spectral sharpness.

Similarly to other machine learning areas, separate training and test sets were defined, consisting of 6 speakers each. The feature selection process including the application of PCA was performed on the training set, as well as the parameter setting of the clustering algorithms (c for the FCM and its variants and α for s-FCM). Then, the test set was transformed in a similar way to the training one (i.e. using the same (basic) features, then transformed by PCA using the same principal components, and keeping the same number of transformed attributes). Lastly, the transformed test set was clustered using the clustering parameter values obtained on the training set, and the result was evaluated using the purity and entropy metrics.

For the pre-processing steps, we experimented with keeping the 20, 50 and 100 most correlated features; after PCA, we kept 75%, 90%, 95% and 99% of the information.

4.4 Parameter Setting for the Clustering Methods

Although $m = 2$ is the most frequently employed value for the fuzzy exponent, it is not suitable when the number of dimensions is several dozens because it leads all cluster prototypes to the grand mean of the input data. In all algorithms that contain the probabilistic exponent m , we tested values in the range of 1.05 to 1.5. For all algorithms that use the possibilistic exponent p , we set $p = m$. Possibilistic penalty terms η_i ($i = 1, \dots, c$) were always chosen equal for all clusters, but their value always depended on the actual number of dimensions d . In case of PFCM algorithm, fine results were obtained for $0.6\sqrt{d} \leq \sqrt{\eta_i} \leq 1.25\sqrt{d}$, while FPPPCM performed best for $1.2\sqrt{d} \leq \sqrt{\eta_i} \leq 2\sqrt{d}$. The actual number of dimensions varied between 2 (for 20 features and 75% PCA) and 54 (for 100 features and 99% PCA). For the suppressed FCM algorithm, all suppression rates multiple of 0.1 were considered, but most accurate results were obtained in the range $0.5 \leq \alpha \leq 0.8$.

As HCM has no parameters at all, it required no parameter adjustment. However, as it is not a robust procedure, for this algorithm we performed 100 clusterings for each preprocessing configuration, and averaged out the resulting purity and entropy scores.

5 Results

The resulting purity scores can be seen in Table 1, while the corresponding entropy values are listed in Table 2. The best values for a pre-processing configuration are shown in **bold**. We can see that by increasing the number of

features, the quality of clustering also improves, and it also improves if we retain more information after the PCA step. When using 100 features, however, the difference is quite small between keeping 95% or 99% of the information; on the other side, it is pointless using fewer than 50 features, or keeping only 75% of the information after the PCA step, as the resulting purity scores are pretty low.

Table 1

Purity scores achieved with the different preprocessing configurations and clustering algorithms

Inform. Kept after PCA	Clustering Method	Number of Features					
		20		50		100	
		Train	Test	Train	Test	Train	Test
75%	HCM	68.6%	59.2%	77.7%	59.9%	77.9%	59.1%
	FCM	68.6%	59.4%	80.5%	61.2%	79.7%	60.4%
	s-FCM	69.1%	58.6%	80.8%	61.7%	79.0%	60.9%
	PFCM	72.7%	57.6%	82.1%	62.5%	80.0%	61.4%
	FPPPCM	73.0%	57.3%	82.3%	66.2%	80.8%	66.4%
90%	HCM	77.1%	64.8%	84.2%	69.3%	80.3%	74.2%
	FCM	76.9%	64.3%	85.2%	76.6%	85.5%	77.6%
	s-FCM	77.1%	65.1%	84.9%	75.5%	85.7%	75.0%
	PFCM	77.1%	65.1%	86.0%	76.0%	87.3%	78.1%
	FPPPCM	76.9%	65.1%	86.8%	82.0%	88.9%	80.7%
95%	HCM	73.3%	67.7%	86.0%	71.4%	79.2%	76.3%
	FCM	74.6%	65.9%	86.0%	80.0%	85.7%	78.1%
	s-FCM	76.6%	65.9%	85.2%	78.4%	86.5%	77.1%
	PFCM	75.1%	67.2%	87.0%	78.4%	88.3%	77.1%
	FPPPCM	75.9%	68.5%	87.8%	76.8%	89.6%	83.1%
99%	HCM	73.5%	66.2%	85.2%	75.0%	80.2%	79.5%
	FCM	75.8%	69.5%	86.0%	80.2%	85.5%	79.7%
	s-FCM	75.1%	70.6%	86.2%	76.0%	87.0%	82.0%
	PFCM	76.9%	64.8%	86.2%	76.0%	86.8%	82.3%
	FPPPCM	76.1%	66.2%	86.8%	82.6%	88.8%	83.9%

Regarding the choice of the clustering method, it is clear that HCM performed the worst; the likely reason for this is that it is a stochastic method. There is no great difference among the performances of the other four methods, but generally, FPPPCM performed best both on the training and on the test sets. This seems to indicate that it is not just a method that can be fine-tuned to suit our needs, but the tuned parameter values perform well on another set of examples (e.g. the test set), meaning that the method is a very robust one.

Table 2

Entropy scores achieved with the different preprocessing configurations and clustering algorithms

Inform. Kept after PCA	Clustering Method	Number of Features					
		20		50		100	
		Train	Test	Train	Test	Train	Test
75%	HCM	0.428	0.544	0.327	0.577	0.307	0.611
	FCM	0.421	0.548	0.328	0.565	0.292	0.598
	s-FCM	0.425	0.559	0.327	0.547	0.289	0.599
	PFCM	0.358	0.552	0.312	0.555	0.293	0.579
	FPPPCM	0.359	0.546	0.277	0.510	0.254	0.544
90%	HCM	0.313	0.456	0.265	0.462	0.285	0.450
	FCM	0.309	0.488	0.280	0.416	0.256	0.434
	s-FCM	0.313	0.453	0.256	0.446	0.262	0.438
	PFCM	0.313	0.453	0.256	0.446	0.242	0.422
	FPPPCM	0.312	0.467	0.235	0.467	0.211	0.372
95%	HCM	0.336	0.442	0.232	0.434	0.305	0.407
	FCM	0.315	0.494	0.270	0.397	0.259	0.405
	s-FCM	0.314	0.476	0.248	0.397	0.255	0.371
	PFCM	0.309	0.442	0.248	0.397	0.233	0.407
	FPPPCM	0.317	0.436	0.238	0.385	0.206	0.335
99%	HCM	0.340	0.491	0.240	0.406	0.290	0.371
	FCM	0.313	0.451	0.263	0.380	0.259	0.404
	s-FCM	0.312	0.472	0.223	0.404	0.241	0.361
	PFCM	0.306	0.443	0.230	0.404	0.247	0.347
	FPPPCM	0.313	0.427	0.254	0.327	0.217	0.328

In general, the variations of fuzzy c -means performed somewhat better than the standard algorithm: the latter achieved its best results with 50 features, while the other three methods were able to utilize the extra information stored in the additional features, thus achieving a clustering, which is of a better quality.

Figure 1 shows the purity scores given by the employed set of clustering algorithm in various scenarios. It is evident that FCM's accuracy drops as the fuzzy exponent m grows beyond a critical value situated around 1.3. Among all tested algorithms, FPPPCM performed the best, while the other fuzzy and possibilistic approaches provided results of approximately same quality, but better than the outcome of HCM. FPPPCM even gave purity scores above 0.85, but that setting never coincided with the best performing scenario on the train data set.

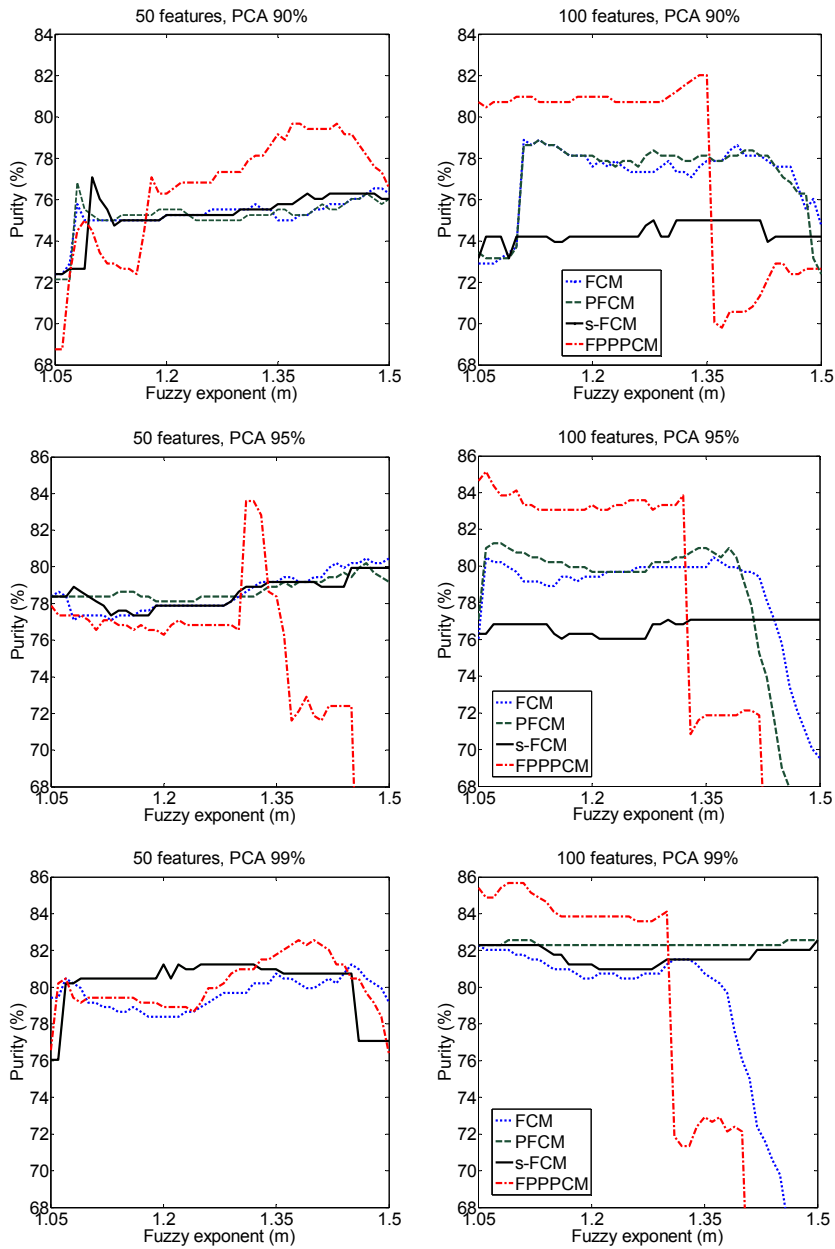


Figure 1

The purity scores obtained plotted against the fuzzy exponent m for the different clustering methods and preprocessing configurations on the training set. In case of 50 features and PCA 99%, the black curve fully covers the green one

The results of two clustering configurations can be seen on Figure 2: the left hand side shows the baseline setting, using all the 6373 features with the standard HCM clustering method, while the right hand side shows the FPPPCM method with the optimal configuration, using 100 features and keeping 99% of the information after PCA. In the latter case, clearly most of the utterances belonging to a given speaker could be mapped in the same cluster (see the rectangles near the diagonal). A number of utterances were assigned to wrong speakers (these form small straight lines). Overall, this clustering is of a much higher quality than the baseline one shown on the left hand side, where some speakers were confused by each other (see the boxes off the diagonal). This is reflected by the accuracy values as well: while the baseline setting had a purity score of 46.9% and an entropy value of 0.560 on the test set, we were able to achieve scores of 83.9% and 0.328, purity and entropy, respectively.

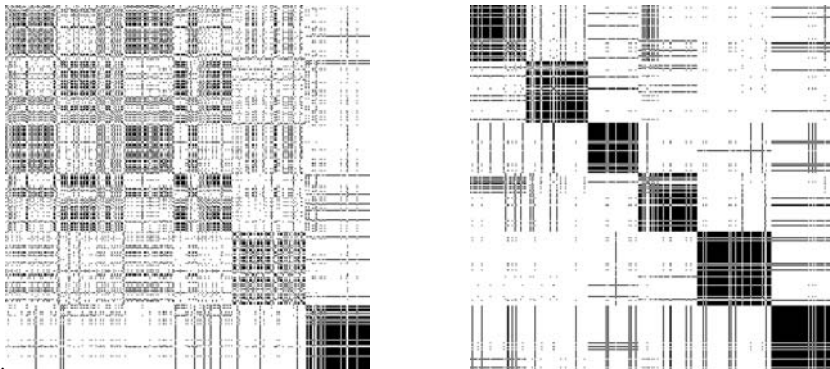


Figure 2

"Confusion matrix" of the test set when using all the features with the HCM method (left), and the best configuration of FPPPCM (right). Each row and column corresponds to one utterance; each point shows whether the corresponding utterances were assigned to the same cluster.

Conclusions

Blind Speaker Clustering (or simply Speaker Clustering) is a task where we have a set of utterances, and our goal is to identify which ones were uttered by the same person. To aid the forming of the desired kinds of clusters, we applied pre-processing steps such as feature selection and Principal Component Analysis (PCA). However, even after these steps it is not trivial to decide which clustering method to apply. Besides the standard algorithms of Hard *c*-means (HCM) and Fuzzy *c*-means (FCM), we tested three further variations of FCM; among them, the fuzzy-possibilistic product partition *c*-means (FPPPCM) proved to be the most effective one, achieving a purity score of 83.9%, which is far above the baseline value of 46.9%.

Acknowledgement

This publication was supported by the Hungarian Government by OTKA under grant no. PD103921.

References

- [1] Rabiner, L., Juang, B. H.: *Fundamentals of Speech Recognition*. Prentice Hall (1993)
- [2] Tóth, L., Tarján, B., Sárosi, G., Mihajlik, P.: Speech Recognition Experiments with Audiobooks. *Acta Cybernetica* 19(4), 695-713 (2010)
- [3] Ondáš, S., Juhár, J., Pleva, M., Lojka, M., Kiktová, E., Sulír, M., Čížmár, A., Holcer, R.: Speech Technologies for Advanced Applications in Service robotics. *Acta Polytechnica Hungarica* 10(5), 45-61 (2013)
- [4] Tóth, S., Sztahó, D., Vicsi, K.: Speech Emotion Perception by Human and Machine. *Proceedings of COST Action*. pp. 213-224, Patras, Greece (2012)
- [5] Gosztolya, G., Grósz, T., Busa-Fekete, R., Tóth, L.: Detecting the Intensity of Cognitive and Physical Load using AdaBoost and Deep Rectifier Neural Networks. *Proceedings of Interspeech*. pp. 452-456, Singapore (Sep 2014)
- [6] Dobry, G., Hecht, R., Avigal, M., Zigel, Y.: Supervector Dimension Reduction for Efficient Speaker Age Estimation based on the Acoustic Speech Signal. *IEEE Trans. Audio, Speech and Language Processing* 19(7), 1975-1985 (2011)
- [7] Räsänen, O., Pohjalainen, J.: Random Subset Feature Selection in Automatic Recognition of Developmental Disorders, Affective States, and Level of Conflict from Speech. *Proceedings of Interspeech*. pp. 210-214, Lyon, France (Sep 2013)
- [8] Gosztolya, G., Busa-Fekete, R., Tóth, L.: Detecting Autism, Emotions and Social Signals using AdaBoost. *Proceedings of Interspeech*. pp. 220-224, Lyon, France (Aug 2013)
- [9] Gosztolya, G.: Conflict Intensity Estimation from Speech using Greedy Forward-Backward Feature Selection. *Proceedings of Interspeech*. Dresden, Germany (2015)
- [10] Schuller, B., Steidl, S., Batliner, A., Epps, J., Eyben, F., Ringeval, F., Marchi, E., Zhang, Y.: The Interspeech 2014 Computational Paralinguistics Challenge: Cognitive & Physical Load. *Proceedings of Interspeech* (2014)
- [11] Gosztolya, G.: Estimating the Level of Conflict Based on Audio Information using INVERSE Distance Weighting. *Acta Universitatis Sapientiae, Electrical and Mechanical Engineering* (2014)
- [12] Kaya, H., Özkaptan, T., Salah, A. A., Gürgen, F.: Canonical Correlation Analysis and Local Fisher Discriminant Analysis-based Multi-View

- Acoustic Feature Reduction for Physical Load Prediction. Proceedings of Interspeech, pp. 442-446, Singapore (Sep 2014)
- [13] van Segbroeck, M., Travadi, R., Vaz, C., Kim, J., Black, M. P., Potamianos, A., Narayanan, S. S.: Classification of Cognitive Load from Speech using an i-vector Framework. Proceedings of Interspeech. pp. 671-675, Singapore (Sep 2014)
- [14] Ajmera, J., Wooters, C.: A Robust Speaker Clustering Algorithm. Proceedings of ASRU. pp. 411-416 (2003)
- [15] Yu, K., Jiang, X., Bunke, H.: Partially Supervised Speaker Clustering. IEEE Transactions on Pattern Analysis and Machine Intelligence 34(5), 959-971 (2012)
- [16] Han, K. J., Narayanan, S. S.: Agglomerative Hierarchical Speaker Clustering using Incremental Gaussian Mixture Cluster Modeling. Proceedings of Interspeech. pp. 20-23 (2008)
- [17] Steinhaus, H.: Sur la division des corps materiels en parties. Bull. Acad. Pol. Sci. C1 III. (IV), 801-804 (1956)
- [18] Bezdek, J. C.: Pattern Recognition with Fuzzy Objective Function Algorithms, Plenum, New York (1981)
- [19] Pal, N. R., Pal, K., Keller, J. M., Bezdek, J. C.: A Possibilistic Fuzzy c -means Clustering Algorithm. IEEE Trans. Fuzzy Syst. 13, 517-530 (2005)
- [20] Szilágyi, L., Szilágyi, S. M.: Generalization Rules for the Suppressed Fuzzy c -means Clustering Algorithm. Neurocomputing 139, 298-309 (2014)
- [21] Devijver, P., Kittler, J.: Pattern Recognition, a Statistical Approach. Prentice Hall (1982)
- [22] Brendel, M., Zaccarelli, R., Devillers, L.: A Quick Sequential Forward Floating Feature Selection Algorithm for Emotion Detection from Speech. Proceedings of Interspeech. pp. 1157-1160, Makuhari, Japan (2010)
- [23] Jolliffe, I.: Principal Component Analysis. Springer-Verlag (1986)
- [24] Busa-Fekete, R., Kocsor, A.: Locally Linear Embedding and its Variants for Feature Extraction. Proceedings of SOFA. pp. 216-222 (2005)
- [25] Krishnapuram, R., Keller, J. M.: A Possibilistic Approach to Clustering. IEEE Trans. Fuzzy Syst. 1, 98-110 (1993)
- [26] Fan, J. L., Zhen, Z. W., Xie, W. X.: Suppressed Fuzzy c -means Clustering Algorithm. Patt. Recogn. Lett. 24, 1607-1612 (2003)
- [27] Szilágyi, L.: Fuzzy-Possibilistic Product Partition: a Novel Robust Approach to c -means Clustering. Proceedings of MDAI. pp. 150-161, Changsha, China (Jul 2011)

- [28] Manning, C., Raghavan, P., Schütze, H.: Introduction to Information Retrieval. Cambridge University Press (2008)
- [29] Todd, S. C., Tóth, M. T., Busa-Fekete, R.: A Matlab Program for Cluster Analysis using Graph Theory. *Computers & Geosciences* 36(6), 1205-1213 (2009)
- [30] Endo, Y., Kinoshita, N., Iwakura, K., Hamasuna, Y.: Hard and Fuzzy *c*-means Algorithms with Pairwise Constraints by Non-Metric Terms. *Proceedings of MDAI*. pp. 145-157, Tokyo, Japan (Oct 2014)
- [31] Schuller, B., Friedmann, F., Eyben, F.: The Munich Biovoice Corpus: Effects of Physical Exercising, Heart Rate, and Skin Conductance on Human Speech Production. *Proceedings of LREC*. pp. 1506-1510, Reykjavik, Iceland (2014)
- [32] Eyben, F., Wöllmer, M., Schuller, B.: Opensmile: The Munich Versatile and Fast Open-Source Audio Feature Extractor. *Proceedings of ACM Multimedia*. pp. 1459-1462 (2010)

Challenges in Preserving Intent Comprehensibility in Software

**Valentino Vranić¹, Jaroslav Porubän², Michal Bystrický¹,
Tomáš Frtála¹, Ivan Polášek^{1,3}, Milan Nosál², Ján Lang¹**

¹ Institute of Informatics and Software Engineering, Faculty of Informatics and Information Technologies, Slovak University of Technology in Bratislava, Ilkovičova 2, 84216 Bratislava, Slovakia, vranic@stuba.sk, to-mas.frtala@stuba.sk, michal.bystricky@stuba.sk, jan.lang@stuba.sk

² Department of Computers and Informatics, Faculty of Electrical Engineering and Informatics, Technical University of Košice, Letná 9, 04200 Košice, Slovakia, jaroslav.poruban@tuke.sk, milan.nosal@tuke.sk

³ Gratex International, a. s., Galvaniho 17/C, 821 04 Bratislava, Slovakia, ipo@gratex.com

Abstract: Software is not only difficult to create, but it is also difficult to understand. Even the authors themselves in a relatively short time become unable to readily interpret their own code and to explain what intent they have followed by it. Software is being created with the goal to satisfy the needs of a customer or directly of the end users. Out of these needs comes the intent, which is relatively well understandable to all stakeholders. By using other specialized modeling techniques (typically the UML language) or in the code itself, use cases and other high-level specification and analytical artifacts in common software development almost completely dissolve. Along with dedicated initiatives to improve preserving intent comprehensibility in software, such as literate programming, intentional programming, aspect-oriented programming, or the DCI (Data, Context and Interaction) approach, this issue is a subject of contemporary research in the re-revealed area of engaging end users in software development, which has its roots in Alan Kay's vision of a personal computer programmable by end users. From the perspective of the reality of complex software system development, the existing approaches are solving the problem of losing intent comprehensibility only partially by a simplified and limited perception of the intent and do this only at the code level. This paper explores the challenges in preserving the intent comprehensibility in software. The thorough treatment of this problem requires a number of techniques and approaches to be engaged, including preserving use cases in the code, dynamic code structuring, executable intent representation using domain specific languages, advanced UML modularization, 3D rendering of UML, and representation and animation of organizational patterns.

Keywords: intent; use cases; domain specific languages; 3D rendering of UML; organizational patterns

1 Introduction

Software is not only difficult to create, but it is also difficult to understand. Even the authors themselves in a relatively short time become unable to readily interpret their own code and to explain what intent they have followed by it, i.e., what they wanted to achieve. Similar situations arise with models and in particular with more detailed, design models. This problem is being solved by introducing another artifact into software development: documentation. This brings in a further complex problem: the need to keep documentation up to date. In the case of internal documentation (comments), the traceability of the artifacts the documentation is related to has to be ensured, too. Furthermore, a considerable effort is needed to initially create the documentation, inevitably with a disputable and difficult to control quality because its consistency, as opposed to that of a program, cannot be tested by actual execution.

This problem can also be perceived in a more global manner. Software is being created with the goal to satisfy the needs of a customer or directly the needs of the end users. Out of these needs comes the intent, which is relatively well understandable to all the stakeholders. By using other specialized modeling techniques (typically the UML language) or in the code itself, use cases and other high-level specification and analytical artifacts in common software development almost completely dissolve.

Understanding the intent expressed in code has been identified as one of the key problems in software development that has a direct impact on creating programming languages and related tools [36]. This is important not only in software maintenance, but it is also related to the question of reuse: the comprehension of the intent as realized by a given component is necessary for its reuse.

Along with dedicated initiatives to improve preserving intent comprehensibility in software, such as literate programming, which subordinates code to documentation [34], intentional programming, which aimed at enabling direct creation of appropriate abstractions by the programmer [57], aspect-oriented programming, which makes possible to gather parts of the code into modules by use cases [25, 26], or the DCI (Data, Context and Interaction) approach, which enables to partially preserve use cases [14], this issue is a subject of the contemporary research in the re-revealed area of engaging end users in software development [6] that has its roots in Alan Kay's vision of a personal computer programmable by end users. From the perspective of the reality of complex software system development, the existing approaches solve the problem of losing intent comprehensibility only partially, by a simplified and limited perception of the intent and do so only at the code level.

This paper explores the challenges in preserving intent comprehensibility in software. The thorough treatment of this problem requires a number of techniques and approaches to be engaged. Section 2 explains the dimensions of the intent in soft-

ware. Sections 3-5 explore the possibilities of preserving intent comprehensibility from the perspective of each of the basic software constituents. Section 6 discusses the related work. The paper is closed by conclusions and indication of further work directions.

2 Dimensions of Intent

The ultimate form of the software is the executable *code*, which is mostly text based. The level of the comprehensibility of the intent expressed by the code determines the quality of the resulting software system: better intent comprehensibility simplifies error discovery and lessens the divergence from the functionality that the software system should have provided.

Software development is accompanied by the creation of many artifacts that as such do not contribute to its functionality and thus fairly quickly become outdated. These additional artifacts are usually conjointly referred to as documentation. A special position among these is held by *models* as non-code artifacts predominately expressed in a graphical form and used to reason about the software system being developed. As such, they can be perceived as a transient form towards the code, which is, after the code has been created, condemned to outdated. Models can also be used to generate code or other models, or they can even be executable. In any case, it is important for the intent in models to be comprehensible, too.

The originators of the intent are *people* and losing the intent in organizing people is transferred to the software being developed by these people. This is a direct consequence of Conway's law [11]:

Organizations which design systems are constrained to produce designs which are copies of the communication structures of these organizations.

This is not only the problem of the initial organization of the people. The people remain the key factor throughout software development including the maintenance phase, too. They tend to literally impersonate the software system parts they develop and the effectiveness of resolving development problems depends directly on the effectiveness of the communication among the developers. This is where agile and lean approaches, which favor face-to-face contact among people over any kind of formal communication, save significant time and resources [14].

Furthermore, the notion of intent in software is relative and it can be observed in the following dimensions:

- *Stakeholders* (intent originators): from whose perspective the intent is observed, starting with the end user and moving towards the programmer

- *Level*: at what level of construct granularity is the intent expressed, starting with the lower level constructs (e.g., conditional statements or loops), via covering constructs (e.g., methods, classes, etc.), conceptual constructs and software system parts (different levels of subsystems), up to the overall software system, including people organization
- *Expression*: how is the intent expressed, starting with an idea or mental model, via informal notes and further forms of textual description, including requirements lists and use cases, up to graphical model representations, formal specification, and, finally, the executable form itself

The entire intent space as determined by these dimensions, i.e., *stakeholders–level–expression*, is huge. With respect to the basic software constituents, three particular areas of interest can be identified therein (see Figure 1):

- With respect to code, it is reasonable to observe the intent by its representation in an executable form at the code level up to the covering constructs level
- With respect to models, the focus should be on expressing the intent by graphical models spanning from the conceptual construct level up to the software system level
- With respect to people, the most interesting is the people organization as such and people organization in projects, at which the employment of textual description and graphical models to express the intent appropriately should be explored

All three areas span throughout the whole stakeholder dimension since, in general, any stakeholder can be involved in any software artifact. This is at heart of the agile and lean approaches to software development with their concept of cross-functional roles. It may seem strange for end users to be connected with code, but it a huge number of people in the USA (four times the number of professional programmers there) reported they do programming at work [6]. With techniques that shape code according to use cases and domain specific languages, addressed in the next section, end user programming becomes even more relevant. Sections 4 and 5 address the remaining two areas of interest in exploring the intent in software.

3 Code Perspective

As we will see in this section, the intent expressed in use cases can actually survive in code (Section 3.1) with application domain abstractions supported by domain specific languages (Section 3.2), while the differences in the mental models of individual stakeholders require abandoning the fixed code structure (Section 3.3).

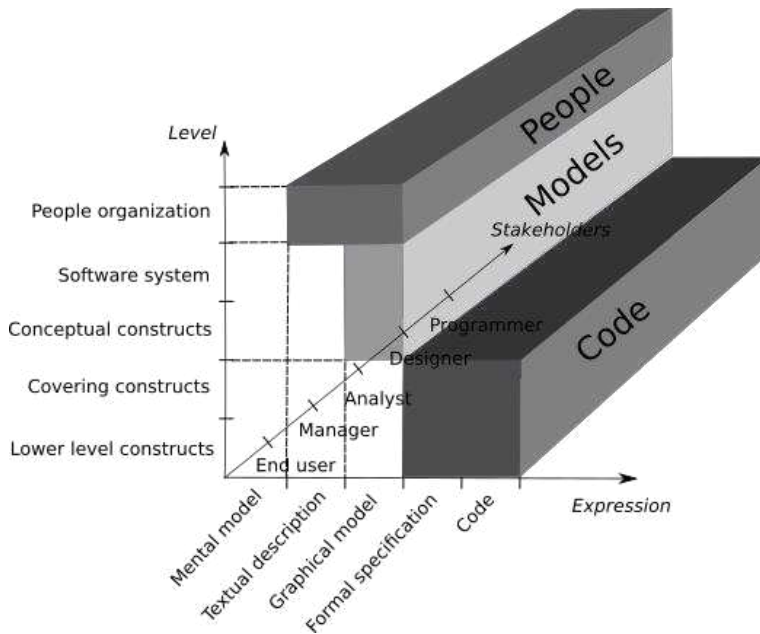


Figure 1

Areas of interest in exploring the intent in software

3.1 Preserving Use Cases

From the perspective of preserving use cases in code, there are two approaches of particular interest: aspect-oriented programming, which enables to collect the code parts into modules corresponding to use cases [25, 26], and the DCI approach [14], which enables to partially preserve use case flows, i.e., sequences of steps in use cases—known also as *flows of events* or simply *flows*¹—albeit they remain fragmented by roles.

In common object-oriented programming, the client or end user intent diminishes from code. The parts of use cases end up in different classes by which they are realized. The ability to affect one use case by another one without having any reference to the affecting use case in the affected one, known as the extend relationship, also lacks in common object-oriented programming.

What should be explored is how appropriate design patterns, the frameworks based on these design patterns, and preprocessing techniques can ensure not only

¹ Sometimes *scenario* is used to denote a use case flow. This may be confusing since a scenario can stand for a particular path through a use case, which involves some or all steps of one or several use case flows.

the code to be modularized according to use cases, but also how use case flows (the actual steps) can be preserved in the code and how to achieve this in a form close to the natural language. The modification of use cases, which usually happens only in code without reflecting the changes in the use case model, would consequently be readable to stakeholders that have no programming knowledge. This would even open a possibility for these stakeholders to directly modify the program, at least at the highest level.

It is necessary to investigate the possibilities of expressing individual steps in use case flows and the possibility of their formal interpretation without fully formalizing how they are expressed. For this, formal processing of informal meaning by abstract interpretation could be considered [35]. By now, it has been demonstrated, though only in the Python programming language, how use case flows can be preserved in code [7]. The modularization of use cases using design patterns in the PHP programming language has been addressed, too [23], but this approach does not preserve use case flows.

3.2 Domain Specific Languages

Domain specific languages are being created with the goal of bringing closer the way the code is expressed to the problem domain. Programmers use problem domain notions directly in the solution, while tools communicate with them in this notional apparatus, too [42]. Domain specific languages play a key role in model driven software development. Because of this, the orientation on domain specific languages in searching for the ways to express the intent in software comes as a natural choice. Despite a growing popularity of domain specific languages, a number of questions remain unanswered in this area. These include language composition [18], effective sentence creation using projectional and hybrid editors [62], and assessment of the usability of domain specific languages [4].

Domain specific languages can be used to achieve a readable and executable intent representation that enables to propagate the notions from use cases, i.e., application domain, to code, i.e., solution domain. In other words, domain specific languages raise the level of abstraction of the solution domain closer to the application domain abstractions making it possible for programmers to express the solution using in the application domain fashion.

There are three promising directions in the research of preserving intent comprehensibility with domain specific languages. The first one is to come closer to the ideal case in which a domain specific language will be both the application domain language and solution domain language. A domain specific language can have textual, but also graphical (visual) concrete representations, which constitutes the grounds for the second direction of research in preserving the intent with domain specific languages: overcoming the gap between the model and the code, while retaining executability. The third direction aims at simplifying the creation

of domain specific languages and their evolution, which takes place along with the evolution of the program itself (constituting a program–language coevolution).

3.3 Dynamic Code Structure

A huge impediment to observing the intent is the fixed code structure. Different stakeholders need to see code in different ways in which—from their perspective—the intent is readable. However, these cannot be based on a static code representation, but have to be modifiable as though they are the actual code representation. The challenge here is not only to design the necessary views, but also to enable creating further views directly by stakeholders.

The fixed code structure is a consequence of the economic aspects of software development. Developers are bound to choose exactly one code structure. This code structure corresponds to the mental model they have built upon their experience and knowledge, but also to the nature of the intent they have to realize. An alternative implementation that would employ some other code structure is in this case redundant and thereof economically inconvenient. Thus, the final code structure usually favors one intent that the authors considered to be the most important. That intent can be anything, including technical intents such as software efficiency, software extensibility, etc.

The programmers that join a project during the course of its realization, have to understand the existing code before they can manage to progress. In such a situation, the new programmers, as new stakeholders, have to adopt the mental model of the original programmers. This is difficult, since their own experience, knowledge, and preferences in general are only rarely close to those of the original programmers and thus constitute a mental barrier to the code comprehension. If, for example, the original programmers focused on the system efficiency, while the new programmers prefer extensibility, they might not understand many of the decisions made by the original programmers making it difficult for them to comply with these decisions.

The problem of the limitations of static structure is in practice partially being attacked by the built-in projections in integrated development environments (IDE). Finding variable uses, which enables to follow scattering of the variable use and thus to follow the intent implemented by it, is one of these. Similarly, environments enable to follow selected intents that are not directly part of the executable code. One example is comments with the *TODO* prefix, which can be found and provided to stakeholders by the IDE in a navigable list with all the instances of a given comment.

Approaches are known that improve certain aspects in the context of the problem area of dynamic code structuring. A prominent example is a method for a faster orientation in code supported by a tool that enables to display the body of the

method being called directly at the place of the call without the need of explicit navigation [16].

Intentional code views from the perspective of the architectural intents based on logical metaprogramming have been reported [41]. However, these views are not editable. Intents have been represented in a form of a graph abstraction, too [55]. Programming with so-called ghosts [8] is based on automatic creation of undefined yet entities used in the program, which are displayed in a separate, editable view. The approach has been implemented as a prototype in the form of an Eclipse plugin and as an extension to Smalltalk Pharo. Recording and automating design pattern application treated, for example, by Kajsa [27], can also be perceived as a way of expressing the intent using metadata.

4 Model Perspective

Dynamic code structuring as addressed in Section 3.2 is conceptually applicable to graphical models, too. Providing different, modifiable views instead of only one, static view is highly related to aspect-oriented modularization or to what is known as advanced modularization in general. There are some opportunities to achieve this in UML as a de facto standard in software modeling (Section 4.1). Employing the third dimension in representing software models can further improve intent comprehensibility (Section 4.2).

4.1 Advanced Modularization in UML

Despite extensive research in the area of aspect-oriented software development, expressing advanced modularization at the model level did not end up with a generally accepted approach. A very important approach in this direction is Theme [10], which is, similarly as newer approaches such as RAM [33], based on non-UML elements.

Separation of concerns with clearly expressing their interrelation naturally contributes to intent comprehensibility. The UML diagrams typically used in practice make this possible only to a limited extent. Therefore, it is necessary to investigate how advanced elements of UML, which usually have no straightforward counterparts in code, can help in expressing intent. In this sense, composite structure models, whose important elements are roles and their collaborations, are particularly interesting. Here, it is necessary to search for the ways of expressing roles and their composition. Another UML concept that is directly interconnected with composite structure models are parameterized types.

4.2 3D Rendering of UML

Model rendering itself and creation of alternative views can also enhance intent comprehensibility by reducing its fragmentedness. The third dimension can be employed here to simultaneously display and interconnect related parts of the model. For this, a complex 3D UML rendering support for the layout of class or module layers and their relationships has to be provided. This is different than the standard package modularization, in which the relationships between package elements are not easily observed. The point is in maintaining the layers in a simulated 3D space in which it will be possible to create and observe elements and their relationships along with the relationships between the layers as such.

In this sense, the model could be structured according to use cases as intent bearers. Use cases define interaction and therefore are commonly modeled by sequence diagrams, which implicitly uncover the underlying structure necessary for the realization of use cases [26, 1]. To expose the structure directly, sequence diagrams can be easily converted into communication or object diagrams and displayed in their own layers. The class diagram automatically created out of the communication or object diagrams could be displayed in another layer making the correspondence—and their intent—of the elements of these different diagrams obvious provided their planar coordinates are preserved. The idea itself has been indicated earlier [47]. Current research efforts aim at realizing it [22]. This approach is applicable also to decoupling patterns and antipatterns in class diagrams depicted schematically in Figure 2.

Several approaches to the 3D rendering of UML diagrams have been reported. However, each of these approaches is targeting only one type of UML diagrams and none of them supports editable 3D views. X3D-UML [39] is an approach to the 3D rendering of UML state machine diagrams in movable hierarchical layers with the possibility of applying filtering. No appropriate tool for editing this view is available.

GEF3D [17, 45] is a 3D framework based on Eclipse GEF (Graphical Editing Framework) developed as an Eclipse plugin. By using this framework, existing GEF-based 2D editors can easily be embedded into 3D editors. The main idea of this framework is to use the third dimension to visualize connections among several common (2D) UML diagrams each of which is displayed in a separate layer parallel to other layers. GEF3D supports also orthogonal positioning of layers into virtual boxes [17] that makes inter-model connections clearly visible, but limits the number of layers. GEF3D views are non-editable. Moreover, the project has not been maintained since 2011.

A different way of 3D rendering of UML diagrams is based on so-called geons [9], simple geometrical forms by which humans recognize more complex objects according to Biederman's recognition-by-components theory. In UML diagrams, a different geon is assigned to each kind of model element. According to this ap-

proach, by getting used to this mapping, even complex diagram structures become readily comprehensible.

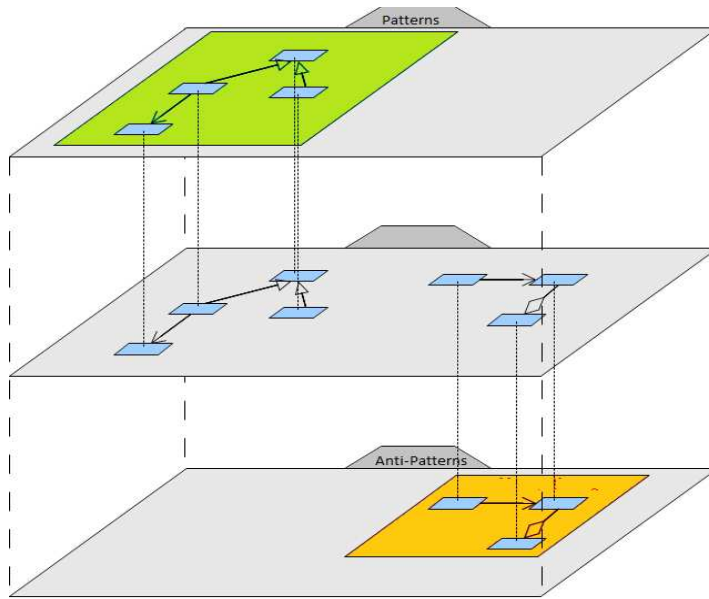


Figure 2

Decoupling patterns and anti-patterns with a layered 3D rendering of a class diagram

5 People Perspective

Appropriate ways of organizing people have been discovered in successful projects of software development and captured as organizational patterns [15]. Approaches that will enable to better understand the intent of organizational patterns individually and in combination have to be explored. One possibility is their clarification using software modeling and UML in particular, including the 3D rendering of UML discussed in Section 3.3. This approach is applicable to expressing the organization of people in areas other than software development, too.

Alternatively, organizational patterns could be modeled in a virtual world in which it would be possible for people to try the roles featured in these patterns and experience the characteristic problem situations in a simulated environment.

For example, in the Architect also Implements pattern [15], a stakeholder could play the role of a programmer who has to understand the design of a software system in order to be able to implement it. The stakeholder could also play the role of a software architect who prepares the design without a clear idea of its

implementation. The stakeholder could also experience each of these roles in a positive arrangement in which the architect cooperates with programmers and contributes to the implementation.

A virtual world does not necessarily mean virtual reality. That would surely be a benefit, but the human imagination is capable of substituting this dimension when important features of the content are captured. This phenomenon is known in videogames, among which those with an interesting story tend to endure despite a simpler graphical workout. Thus, the essence is in creating the corresponding model of a typical situation solved by a given organizational pattern. In general, it is necessary to find an approach of transforming organizational patterns into such typical situations and to create a framework in which they could be readily expressed.

Agile games [37, 20] provide a possibility to experience situations in which it is possible to better understand principles, relationships, and forces acting in agile and lean approach to software development. The same approach could be applied with organizational patterns.

Differently than agile games, the envisaged representation and animation of organizational patterns in a virtual world would provide a more realistic picture of the situation. It would also take less time and would be applicable in an individual setting with no need to engage other people, nor depend on their time. A simplified representation of this idea is offered by the SimSYS environment [12]. Animating organizational patterns as text adventure games [21] promises to improve the comprehensibility of original descriptions of organizational patterns [15].

6 Related Work

The problem of preserving intent comprehensibility accompanies software development from its beginnings. In the introduction, we indicated some important historical points starting from Alan Kay's vision of personal computer programmable by end users [32] through several approaches to preserving intent in programming, namely literate programming [34], intentional programming [56], aspect-oriented programming [25, 26], and the DCI approach [14]. Preserving intent comprehensibility is a subject of research in contemporary area of engaging end users in software development [6]. Even though involving the client or end users in software development is important, too, we claim that the successful treatment of the problem of preserving intent comprehensibility has to comprise all three basic software constituents: code, models, and people.

Aspect-oriented change realization [5, 40, 64, 65] enables modular expression of a change, by which it actually contributes to a more comprehensible representation of the intent. Determining the realization type of a change that has to be applied is

possible also by the multi-paradigm design with feature modeling [40], which is in its own right interesting from the perspective of preserving intent comprehensibility because it represents an effort to bridge the gap between the solution and application domain by a transformation [63].

The intent in code can be exposed by encouraging programmers to record their intent in the form of intentional comments prior to writing any code as in the design intention driven programming [38]. Such enforcement of documentation directly in the code addresses some of problems and limitations of literate programming [58].

Approaches strongly bound to a model, e.g., domain driven design [19], do not consider exposing behavior at a higher, use case level, but rather they encourage programmers to create knowledge-rich models—also called smart models—which do not evolve well [13]. As a result, the higher-level use cases become fragmented in models and they are not visible in the code. This is addressed by several approaches that preserve the intent in code at a higher level by the modularization of the code into use cases [7, 14, 25, 26].

A use case modeling metamodel can serve as a basis for reasoning about preserving use cases in the code and models [66, 67].

In applying advanced modularization for the purpose of a clearer separation of the intent in code, techniques of applying advanced modularization in established programming languages that are not being denoted as aspect-oriented [3] are of particular interest.

The problem of effective design of domain specific languages from the perspective of tool support [52] is the key to a successful adoption of domain specific languages in software development. A domain specific language is a dynamic element in software development undergoing constant changes. This evolution of domain specific languages may involve their composition. In building domain specific languages, creating their concrete syntax out of the samples of sentences of abstract conception [53], as well as inferring domain specific languages out of user interfaces of existing software systems [1], can help significantly.

Recording the applied design patterns using annotations [56] can help in preserving intent comprehensibility in code. The method of abstracting information from the details of the format being used, targeting XML formats and annotations [43], can be applied here. Furthermore, an analysis of the need for dynamic code structure and its conceptual design has been reported [44], supported by a NetBeans module prototype that enables simple intent based code projections expressed by structured comments [54].

Using 3D space for software modeling in UML has been reported to support code refactoring and optimization by displaying patterns in a separate layer, as well as antipatterns for refactoring in another layer [46, 48, 50, 60]. Code visualization upon AST project graphs [49, 51, 61] can be used to present models. This is relat-

ed not only to algorithms of positioning elements and their relationships, such as FM3 or Fruchterman-Reingold, but also to the internal conception of presenting information in code, such as authors, antipatterns, types of classes, and similar [22].

To successfully master the graphical demands while working with UML models in a 3D space, advanced software visualization [24, 28, 30, 59] and graph visualization in general [31], as well as design of environments for visual programming [29], are necessary.

Conclusions

Up to now, preserving intent comprehensibility has been approached to in a fragmented manner without clearly understanding its relativity. We envisage an integral perception of the intent in software and in support to its comprehensibility in the whole space formed by the identified dimensions of the intent, i.e., stakeholders–level–expression, and in all three basic software constituents, i.e., code, models, and people. Our hypothesis is that it is possible to increase the comprehensibility of the intent by applying the corresponding methods conceived with respect to the dimensions in which the intent is observed:

- Having use cases as part of the code can overcome current fragmentedness of the representation of use case flows, as well as readability of their steps as such
- Domain specific languages can provide a readable and executable intent representation
- Dynamic code structuring brings in editable adaptable code views
- Advanced modularization in UML strives at employing standard yet underutilized UML elements to express the intent
- 3D rendering of UML has the potential of providing a fully editable model capable of displaying the relationship of use cases to the corresponding detailed UML model (applicable to the DCI approach, too), but also patterns and antipatterns in optimization and refactoring, aspects in aspect-oriented modeling, and alternative and parallel scenarios
- Representation and animation of organizational patterns of software development will make possible to transfer the experience of proven ways of organizing people in software development in a new form available to all stakeholders on an individual basis and at a convenient time

Acknowledgement

The work reported here was supported by the Scientific Grant Agency of Slovak Republic (VEGA) under the grant No. VG 1/1221/12.

This contribution/publication is also a partial result of the Research & Development Operational Programme for the project Research of Methods for Acquisition, Analysis and Personalized Conveying of Information and Knowledge, ITMS 26240220039, co-funded by the ERDF.

References

- [1] J. Arlow and I. Neustadt. UML 2 and the Unified Process: Practical Object-Oriented Analysis and Design. 2nd Edition, Addison-Wesley, 2005
- [2] M. Bačíková, J. Porubán, D. Lakatoš: Defining Domain Language of Graphical User Interfaces. In Proceedings of SLATE '13, 2nd Symposium on Languages, Applications and Technologies, Porto, Portugal, 2013
- [3] J. Bálík and Valentino Vranić. Symmetric Aspect-Orientation: Some Practical Consequences. In Proceedings of NEMARA 2012: International Workshop on Next Generation Modularity Approaches for Requirements and Architecture, at AOSD 2012, Potsdam, Germany, ACM, 2012
- [4] A. Barišić, V. Amaral, M. Goulão, and B. Barroca. Quality in Use of Domain-Specific Languages: A Case Study. In Proceedings of 3rd ACM SIGPLAN Workshop on Evaluation and Usability of Programming Languages and Tools, ACM, 2011
- [5] M. Bebjak, V. Vranić, and P. Dolog. Evolution of Web Applications with Aspect-Oriented Design Patterns. In Proceedings of ICWE 2007 Workshops, 2nd International Workshop on Adaptation and Evolution in Web Systems Engineering, AEWSE 2007, in conjunction with ICWE 2007, Como, Italy, 2007
- [6] M. M. Burnett and B. A. Myers. 2014. Future of End-User Software Engineering: Beyond the Silos. In Proceedings of the Future of Software Engineering, FOSE 2014, 36th International Conference on Software Engineering, ICSE 2014, Hyderabad, India, ACM, 2014
- [7] M. Bystrický and V. Vranić. Preserving Use Case Flows in Source Code. In Proceedings of 4th Eastern European Regional Conference on the Engineering of Computer Based Systems, ECBS-EERC 2015, Brno, Czech Republic, IEEE Computer Society, 2015
- [8] O. Callau and É. Tanter. Programming with Ghosts. IEEE Software, 30(1):74-80, 2013
- [9] K. Casey and C. Exton. A Java 3D Implementation of a Geon Based Visualisation Tool for UML. In Proceedings of the 2nd International Conference on Principles and Practice of Programming in Java, Kilkenny City, Ireland, 2003
- [10] S. Clarke and E. Baniassad. Aspect-Oriented Analysis and Design: The Theme Approach Addison-Wesley, 2005

- [11] M. E. Conway. How do Committees Invent?, *Datamation*, 14(4):28-31, April 1968
- [12] K. Cooper and C. Longstreet. Towards Model-Driven Game Engineering for Serious Educational Games: Tailored Use Cases for Game Requirements. In *Proceedings of 17th International Conference on Computer Games, CGAMES*, IEEE, 2012
- [13] J. O. Coplien. The DCI Architecture: Supporting the Agile Agenda. *Øredev Developer Conference*, 2009. <https://vimeo.com/8235574>
- [14] J. O. Coplien and G. Bjørnvig. *Lean Architecture: for Agile Software Development*. Wiley, 2010
- [15] J. O. Coplien and Neil B. Harrison. *Organizational Patterns of Agile Software Development*. Prentice Hall, 2004
- [16] M. Desmond, M.-A. Storey, and C. Exton. Fluid Source Code Views. In *Proceedings of 14th IEEE International Conference on Program Comprehension, ICPC'06*, Washington, DC, USA, IEEE Computer Society, 2006
- [17] K. Duske. A Graphical Editor for the GMF Mapping Model. *GEF3D Development Blog*, 2010. <http://gef3d.blogspot.sk/2010/01/graphical-editor-for-gmf-mapping-model.html>
- [18] S. Erdweg, P. G. Giarrusso, and T. Rendel. Language Composition Untangled. In *Proceedings of 12th Workshop on Language Descriptions, Tools, and Applications*, ser. *LDTA '12*, New York, NY, USA: ACM, 2012
- [19] E. Evans. *Domain-Driven Design: Tackling Complexity in the Heart of Software*. Addison Wesley, 2003. ISBN: 0-321-12521-5
- [20] J. M. Fernandes, S. M. Sousa. PlayScrum—A Card Game to Learn the Scrum Agile Method. In *Proceedings of 2nd International Conference on Games and Virtual Worlds for Serious Applications*, 2010
- [21] T. Frt'ala and V. Vranić. Animating Organizational Patterns. In *Proceedings of 8th International Workshop on Cooperative and Human Aspects of Software Engineering, CHASE 2015, ICSE 2015 Workshop*, Florence, Italy, IEEE, 2015
- [22] L. Gregorovič, I. Polášek, and B. Sobota. Software Model Creation with Multidimensional UML. In *Proceedings of International Conference on Research and Practical Issues of Enterprise Information Systems, CONFENIS 2015, 23rd IFIP World Computer Congress International Conference on Research and Practical Issues of Enterprise Information Systems, LNCS*, Daejeon, Korea, Springer, 2015 (accepted)
- [23] J. Greppel and V. Vranić. An Opportunistic Approach to Retaining Use Cases in Object-Oriented Source Code. In *Proceedings of 4th Eastern European Regional Conference on the Engineering of Computer Based Systems, ECBS-EERC 2015*, Brno, Czech Republic, IEEE Computer Society, 2015

- [24] F. Grznár and P. Kapec. Visualizing Dynamics of Object Oriented Programs with Time Context. In Proceedings of Spring Conference on Computer Graphics, SCCG 2013, Smolenice, Slovakia, ACM, 2013
- [25] I. Jacobson. Use Cases and Aspects—Working Seamlessly Together. *Journal of Object Technology*, 2(4):7–28, 2003
- [26] I. Jacobson and P.-W. Ng. *Aspect-Oriented Software Development with Use Cases* Addison-Wesley, 2004
- [27] P. Kajsa and P. Návrát. Design Pattern Support Based on the Source Code Annotations and Feature Models. In Proceedings of 38th International Conference on Current Trends in Theory and Practice of Computer Science, SOFSEM'12, Springer, 2012
- [28] P. Kapec. Hypergraph-Based Software Visualization. In Proceedings of International Workshop on Computer Graphics, Computer Vision and Mathematics, GraVisMa 2009, Plzeň, Czech Republic, 2009
- [29] P. Kapec. Visual Programming Environment Based on Hypergraph Representations. In Proceedings (Part II) of ICCVG 2010, International Conference on Computer Vision and Graphics, Warsaw, Poland, LNCS 6375, Springer, 2010
- [30] P. Kapec. Visualizing Software Artifacts Using Hypergraphs. In Proceedings of Spring Conference on Computer Graphics in Cooperation, SCCG'2010, Budmerice, ACM, 2010
- [31] P. Kapec, Michal Paprčka, Adam Pažitnaj, and Vladimír Polák. Exploring 3D GPU-Accelerated Graph Visualization with Time-Traveling Virtual Camera. *Journal of Theoretical and Applied Computer Science (JTACS)*, 7(2):16-30, 2013
- [32] A. Kay. A Personal Computer for Children of All Ages. In Proceedings of the ACM Annual Conference – Volume 1 (ACM '72), ACM, 1972
- [33] J. Kienzle, W. Al Abed, F. Fleurey, J.-M. Jézéquel, and J. Klein. Aspect-Oriented Design with Reusable Aspect Models. *Transactions on Aspect-Oriented Software Development*, 7:279-327, 2010
- [34] D. E. Knuth. Literate Programming. *The Computer Journal*, 27:97–111, 1984. <http://www.literateprogramming.com/knuthweb.pdf>
- [35] J. Kollár. Formal Processing of Informal Meaning by Abstract Interpretation. In Proceedings of Smart Digital Futures 2014, SDF-14, Frontiers in Artificial Intelligence and Applications, Volume 262, Chania, Greece, IOS Press, 2014
- [36] T. D. LaToza and B. A. Myers. Hard-to-Answer Questions About Code. In Proceeding of Workshop on Evaluation and Usability of Programming Languages and Tools, PLATEAU '10, Reno, Nevada USA, ACM, 2010

- [37] T. Lynch et al. An Agile Boot Camp: Using a LEGO®-Based Active Game to Ground Agile Development Principles. In Proceedings of 2011 Frontiers in Education Conference, FIE 2011, Rapid City, SD, USA, 2011
- [38] K. Matz. Designing and Evaluating an Intention-Based Comment Enforcement Scheme for Java. MA thesis. Walton Hall, Milton Keynes, MK7 6AA, United Kingdom: The Open University, 2010
- [39] P. McIntosh. X3D-UML: User-Centred Design. Implementation and Evaluation of 3D UML Using X3D. PhD thesis, RMIT University, 2009
- [40] R. Menkyna and V. Vranić. Aspect-Oriented Change Realization Based on Multi-Paradigm Design with Feature Modeling. In Proceedings of 4th IFIP TC2 Central and East European Conference on Software Engineering Techniques, CEE-SET 2009, Revised Selected Papers, LNCS 7054, Krakow, Poland, Springer, 2012
- [41] K. Mens, B. Poll, and S. González. Using Intentional Source-Code Views to Aid Software Maintenance. In Proceedings of International Conference on Software Maintenance, ICSM '03, Washington, DC, USA, IEEE Computer Society, 2003
- [42] M. Mernik, J. Heering, A. M. Sloane. When and How to Develop Domain-Specific Languages. ACM Computing Surveys (CSUR), 37(4):316-344, 2005
- [43] M. Nosál' and J. Porubán. Supporting Multiple Configuration Sources Using Abstraction. Central European Journal of Computer Science. 2(3):283–299, 2012
- [44] M. Nosál', J. Porubán, and M. Nosál'. Concern-Oriented Source Code Projections In Proceedings of Federated Conference on Computer Science and Information Systems, FEDCSIS 2013, Kraków, Poland, IEEE, 2013
- [45] J. von Pilgrim and K. Duske. GEF3D: a Framework for Two-, Two-and-a-Half- and Three-Dimensional Graphical Editors. Proceedings of 4th ACM Symposium on Software visualization, Ammersee, Germany, 2008
- [46] R. Pipík and I. Polášek. Semi-Automatic Refactoring to Aspect-Oriented Platform. In Proceedings of 14th IEEE International Symposium on Computational Intelligence and Informatics, CINTI 2013, Budapest, IEEE, 2013
- [47] I. Polášek. 3D Model ako spôsob návrhu objektovej štruktúry (3D Model for Object Structure Design). In: Systémová integrace, 11(2):82-89, 2004, In Slovak
- [48] I. Polášek, P. Líška, J. Kelemen, and J. Lang. On Extended Similarity Scoring and Bit-Vector Algorithms for Design Smell Detection. In Proceedings of IEEE 16th International Conference on Intelligent Engineering Systems, INES 2012, Lisbon, Portugal, IEEE, 2012

- [49] I. Polášek, I. Ruttkay-Nedecký, P. Ruttkay-Nedecký, T. Tóth, A. Černík, and P. Dušek. Information and Knowledge Retrieval within Software Projects and their Graphical Representation for Collaborative Programming. *Acta Polytechnica Hungarica*, 10(2):173-192, 2013
- [50] I. Polášek, S. Snopko, and I. Kapustík. Automatic Identification of the Anti-Patterns Using the Rule-Based Approach. In *Proceedings of IEEE 10th Jubilee International Symposium on Intelligent Systems and Informatics, SISY 2012, Subotica, Serbia, IEEE*, 2012
- [51] I. Polášek and M. Uhlár. Extracting, Identifying and Visualisation of the Content, Users and Authors in Software Projects. *Transactions on Computational Science XXI: Special Issue on Innovations in Nature-Inspired Computing and Applications, LNCS 8150, Springer Berlin Heidelberg*, 2013
- [52] J. Porubán, M. Forgáč, M. Sabo, and M. Běhálek: Annotation Based Parser Generator. In: *Computer Science and Information Systems*, 7(2):291–307, 2010
- [53] J. Porubán, J. Kollár, M. Sabo: Abstraction of Computer Language Patterns: The Inference of Textual Notation for a DSL. In: *Formal and Practical Aspects of Domain-Specific Languages: Recent Developments*, IGI Global, 2013
- [54] J. Porubán and M. Nosál. Leveraging Program Comprehension with Concern-Oriented Source Code Projections. In *Proceedings of Slate'14, 3rd Symposium on Languages, Applications and Technologies, Bragança, Portugal*, 2014
- [55] M. P. Robillard and G. C. Murphy. Concern Graphs: Finding and Describing Concerns Using Structural Program Dependencies. In *Proceedings of 24th International Conference on Software Engineering, ICSE '02, Orlando, Florida, USA, ACM*, 2002
- [56] M. Sabo and J. Porubán. Preserving Design Patterns Using Source Code Annotations. *Journal of Computer Science and Control Systems*, 2(1):53–56, 2009
- [57] C. Simonyi. The Death of Computer Languages, The Birth of Intentional Programming. Technical report, MSR-TR-95-52, Microsoft Research, 1995
- [58] M. Smith. Towards Modern Literate Programming. Honours project report, University of Canterbury, New Zealand, 2001
- [59] M. Šperka, P. Kapec, and I. Ruttkay-Nedecký. Exploring and Understanding Software Behaviour Using Interactive 3D Visualization In *Proceedings of 8th International Conference on Emerging eLearning Technologies and Applications, ICETA 2010, Stará Lesná, The High Tatras, Slovakia*, 2010

- [60] M. Štolc and I. Polášek. A Visual Based Framework for the Model Refactoring Techniques. In Proceedings of 8th IEEE International Symposium on Applied Machine Intelligence and Informatics, SAMI 2010, Herľany, Slovakia, IEEE, 2010
- [61] M. Uhlár and I. Polášek. Extracting, Identifying and Visualisation of the Content in Software Projects. In Proceedings of 2012 4th World Congress on Nature and Biologically Inspired Computing, NaBIC 2012, Mexico City, Mexico, IEEE, 2012
- [62] M. Voelter, J. Siegmund, T. Berger, and B. Kolb. Towards User-Friendly Projectional Editors. In Proceedings of 7th International Conference on Software Language Engineering, SLE 2014, 2014
- [63] V. Vranić. Multi-Paradigm Design with Feature Modeling. *Computer Science and Information Systems Journal (ComSIS)*, 2(1):79-102, 2005
- [64] V. Vranić, M. Bebjak, R. Menkyna, and P. Dolog. Developing Applications with Aspect-Oriented Change Realization. In Proceedings of 3rd IFIP TC2 Central and East European Conference on Software Engineering Techniques, CEE-SET 2008, Revised Selected Papers, LNCS 4980, Brno, Czech Republic, Springer, 2011
- [65] V. Vranić, R. Menkyna, M. Bebjak, and P. Dolog. Aspect-Oriented Change Realizations and Their Interaction. *e-Informatica Software Engineering Journal*, 3(1):43-58, 2009
- [66] V. Vranić and L. Zelinka. A Configurable Use Case Modeling Metamodel with Superimposed Variants. *Innovations in Systems and Software Engineering: A NASA Journal*, 9(3):163–177, Springer, 2013
- [67] L. Zelinka and V. Vranić. A Configurable UML Based Use Case Modeling Metamodel. In Proceedings of 1st Eastern European Regional Conference on the Engineering of Computer Based Systems, ECBS-EERC 2009, Novi Sad, Serbia, IEEE Computer Society, 2009

Optimization Algorithms in Function of Binary Character Recognition

Petar Čisar¹, Sanja Maravić Čisar², Dane Subošić¹, Predrag Đikanović³, Slaviša Đukanović³

¹Academy of Criminalistic and Police Studies, Cara Dušana 196, 11080 Zemun, Serbia, petar.cisar@kpa.edu.rs, dane.subosic@kpa.edu.rs

²Subotica Tech, Marka Oreškovića 16, 24000 Subotica, Serbia, sanjam@vts.su.ac.rs

³Republic of Serbia, Ministry of Interior, Kneza Miloša 101, 11000 Belgrade, Serbia, predrag.djikanovic@mup.gov.rs, slavisa.djukanovic@mup.gov.rs

Abstract: The paper gives an analysis of some optimization algorithms in computer sciences and their implementation in solving the problem of binary character recognition. The performance of these algorithms is analyzed using the Optimization Algorithm Toolkit, with emphasis on determining the impact of parameter values. It is shown that these types of algorithms have shown a significant sensitivity as far as the choosing of the parameters was concerned which would have an effect on how the algorithm functions. In this sense, the existence of optimal value was found, as well as extremely unacceptable values.

Keywords: binary character recognition; Optimization Algorithm Toolkit; hill climbing; random search; clonal selection algorithm; parameter tuning

1 Introduction

In different areas of mathematics, statistics, empirical sciences, computer science, mathematical optimization makes up the selection of the prime components, according to given requirements from the series of possible alternatives [1]. An optimization problem involves finding the minimum and maximum of a real function by systematically changing the input values from within a pre-defined set and calculating the corresponding values of the function. Taking a wider approach, it can be stated that optimization comprises the detection of extreme values of a particular function concerning a specific domain, including various types of objective functions, as well as assorted types of domains.

Character recognition is a process by which a computer recognizes different characters (letters, numbers or symbols) and turns them into a digital form that a

computer can use; also known as optical character recognition (OCR). OCR is the mechanical or electronic transformation of scanned or photographed images of written or printed characters into a form that the computer can manipulate. It has been researched in detail and usually focuses on scanned documents. Since the foreground and background are easily segmented, high accuracy has been obtained for binary character recognition. In recent years, more and more OCR requests have come from camera-based and web images. In these images, characters usually have no clear foreground or clean background. Furthermore, degradations such as blurring, low resolution, luminance variance, complicated background, distortion, etc. often appear. All these factors contribute to a great challenge in the recognition of the character of natural scene images.

The majority of character recognition methods managed to improve traditional OCR. Weinman *et al.* [2] proposed a scene text recognition method. His emphasis was to combine lexicon information in order to improve scene text recognition result. Yokobayashi *et al.* [3] utilized a GAT (Global Affine Transformation) correlation method to improve the tolerance of binary character recognition to the distortions of scaling, rotation and shearing. GAT method was much time-consuming. De Campos *et al.* [4] utilized a bag-of-visual-word based method to recognize the character of a given natural scene. Six local descriptor-based character features were compared. Although some of them proved to be preferable to traditional OCR, the results still showed the difficulty in this field.

Besides the OCR applications, one can use binary images with great success with various applications given the lower computational complexity solutions, some embedded systems, autonomous robots, intelligent vehicles, using different algorithms such as hybrid evolutionary algorithm and so on.

The authors of this work will analyze the impact of parameters of several optimization algorithms on accuracy in solving the binary character recognition problem. An analysis of optimization algorithms is performed in the Optimization Algorithm Toolkit (OAT) environment. The OAT is a workbench and toolkit for developing, evaluating and experimenting with classical and state-of-the-art optimization algorithms on standard benchmark problem domains [5]. This open source software contains reference algorithm implementations, graphing, visualizations and other options. The OAT offers a functional computational intelligence library so as to investigate algorithms and problems at hand, while also making use of new problems and algorithms. A plain explorer and experimenter graphical user interface is placed on top of this in order to ensure a fundamental comprehension of the library's functionality. Non-technical access is ensured by the graphical user interface for the configuration and the visualization of existing methods on standard benchmark problem examples [5], [16].

For testing the quality of binary character recognition the following optimization algorithms from the OAT were used: parallel mutation hill climber and random search.

The paper consists of six sections. The introduction offers a terminological basis for the problem of character recognition, then in the second section an overview of hill climbing optimization is presented. The clonal selection algorithm is described in the third section, followed by the fourth section dealing with random search optimization. The results of the accuracy examination of analyzed optimization algorithms in binary character recognition and its various aspects in an adequate software environment are given in the fifth section. Finally, this paper closes with the conclusion based on the analyzed cases.

2 Hill Climbing

Hill climbing in mathematics is an optimization method categorized under local search. Further, hill climbing algorithm is known as the discrete optimization algorithm [6]. This algorithm belongs to the category of heuristic search as it uses a heuristic function to compute the next state. Hill climbing algorithm works on the basis of choosing the nearest neighbor as it computes a state that is better than other succeeding states at its current position. Therefore, it is also treated as a greedy algorithm.

Hill climbing is a suitable local optimum detection (a solution not improvable by considering a neighboring value) but there is no certainty to define the best possible solution (the global optimum) [17]. The characteristics that are provided only by the local optima can be improved by using restarts (repeated local search), or more complex iterative schemes (iterated local search), on memory (reactive search optimization and tabu search), or memory-less stochastic modifications (simulated annealing).

The implementation of this algorithm is rather straightforward, thus ensuring its popularity [17]. It is used widely in the area of artificial intelligence, for reaching a goal state from a starting node. While more advanced algorithms, including the simulated annealing and tabu search, might yield preferable results, with specific cases it is hill climbing that will prove to be a better choice. Hill climbing could frequently bring a better result as opposed to different algorithms in the case when there is only a limited period of time available for performing such a search (which is a common real situation). It is an anytime algorithm: it is able to offer a valid solution even if there is an interruption before the end.

Mathematical description [7] - Hill climbing tries to maximize (or minimize) a function $f(x)$, where x represent discrete states. Such states are usually signified by vertices in a graph, in which edges define nearness or similarity of a graph. Hill climbing traces the graph from vertex to vertex, it will consistently locally increase (or decrease) the value of f , until reaching a local maximum (or local minimum) x_m . Further, hill climbing can operate on a continuous space: should

that be the case, the algorithm is known as the gradient ascent (or gradient descent if the function is minimized).

Drawbacks occurring in the application of this algorithm are [7]:

- **Local maxima** - One of the problems regarding hill climbing is that the scope of detection is limited to local maxima. Assuming that the heuristic is not convex, the global maximum may not be reached. Various different local search algorithms will attempt to circumvent this issue (stochastic hill climbing, random walks and simulated annealing). A possible solution to the current problem of hill climbing is the implementation of random hill climbing search techniques.
- **Ridges** - A ridge can be defined as a curve in the search place leading to a maximum, while the ridge's orientation – in comparison with the available moves that are used to climb – ensures that every move leads to a smaller point. To sum it up, every point on a ridge appears to the algorithm as a local maximum, despite the fact that the given point may form part of a curve which leads to a better optimum.
- **Plateau** - A further problem that sometimes occurs with hill climbing is the issue of a plateau. One is faced with a plateau if the search space is flat - a path where the heuristics are situated in great proximity to each other. In the given cases, it may not be possible for the hill climber to determine the direction of the next step, this might end up wandering in a direction not leading to improvement.

The hill climbing algorithm is not a complete algorithm as it does not always give the result even if it exists [6]. The algorithm may fail if it finds any of these problems. It may choose a wrong path resulting in a dead end.

Hill climbing may be implemented with any problem of choice provided that the current state will allow for an accurate evaluation function. As instances for hill climbing one may list the problem of the traveling salesman, the eight-queen problem, circuit design, and a variety of other real-world problems. Inductive learning models have also implemented hill climbing. One specific instance is PALO, a probabilistic hill climbing system which provides a model of inductive and speed-up learning. Certain implementations of this system have been embedded into “explanation-based learning systems”, and “utility analysis” models [8].

Another field of application of hill climbing is robotics, namely to manage multiple-robot teams. A particular case would be the Parish algorithm, which enables scalable and efficient coordination in multi-robot systems. Their algorithm makes it possible for robots to opt for working alone or in teams by using hill-climbing. Robots executing Parish are thus “collectively hill-climbing according to local progress gradients, but stochastically make lateral or downward moves to help the system escape from local maxima” [9].

The procedure of parallel hill climbing is a hill climbing scheme initiated several instances at a time, this results in multiple hills being claimed in parallel.

It may be argued that the most wide-spread use of the stochastic hill climbing algorithm is by Forrest and Mitchell [10], who posited the Random Mutation Hill Climbing (RMHC) algorithm in a paper analyzing the behavior of the genetic algorithm on a deceptive class of bit-string optimization problems.

3 The Clonal Selection Algorithm¹

Artificial Immune Systems (AIS) are adaptive systems, inspired by theoretical immunology and observed immune functions, principles and models, which are applied to problem solving [11].

As it is emphasized in [16], the clonal selection approach inspired the development of the AIS which executes optimization and pattern recognition problems. This initiates from the antigen controlled maturation of B-cells, with associated hyper mutation process. The B-cell is a type of white blood cell and, more specifically, a type of lymphocyte. These immune systems also implement the concept of memory cells to continue to provide excellent ways to solve the problem in question. De Castro and Timmis [11] focused on two vital characteristics of affinity maturation in B-cells. The first, being that the increase of B-cells is in proportion to the affinity of the antigen binding it. Therefore, the greater the affinity, the greater numbers of clones are created. Furthermore, the mutations resulted through the antibody of a B-cell are inversely proportional to the affinity of the antigen it binds. De Castro and Von Zuben [12] developed and frequently used the clonal selection based-AIS (CLONALG) by implementing these two features, which has been applied for executing pattern matching and multi-modal function optimization tasks.

The general principle of functioning of the clonal selection algorithm is presented in Figure 1 [16].

¹ Section 3 draws on material [16] published in Acta Polytechnica Hungarica. Vol. 11, No. 4, 2014

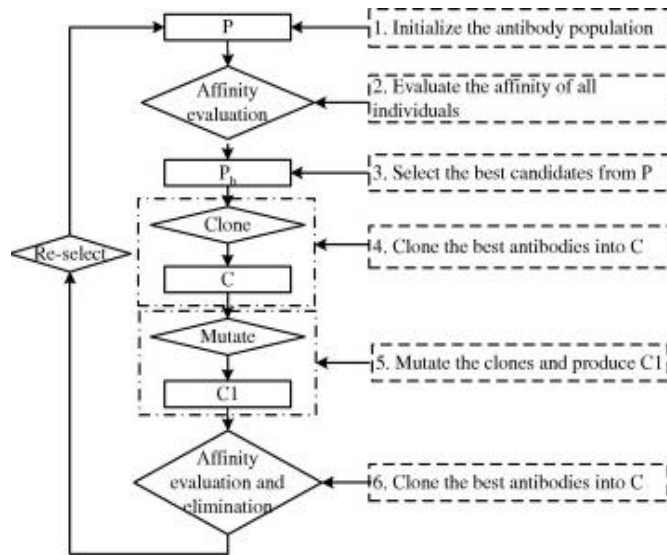


Figure 1
Steps of clonal selection method [13]

4 Random Search

Random Search (RS) represents an optimization method with the aim to detect the optimum by testing the objective function's value on a series of combinations of random values of the adjustable parameters. The random values are generated in compliance with certain boundaries as chosen by the user, moreover, the system excludes any combinations of parameter values failing to fulfill the constraints on the variables. In fact, it can be stated that the method can handle bounds on the adjustable parameters and fulfill constraints. For an infinite number of iterations it is guaranteed that the present method will detect the global optimum of the objective function. Generally, one is interested in processing a very large number of iterations [14]

Options for RS [14]:

- Number of iterations - this parameter is a positive integer to determine the number of parameter sets to be drawn before the algorithm stops. The default value is usually 100000.
- Random number generator - the parameter determines the random number generator, which is to be used with this method.
- Seed - the parameter is a positive integer value with the aim of determining the seed for the random number generator.

The name, random search, is attributed to Rastrigin [15] who made an early presentation of RS along with basic mathematical analysis.

Classification of RS [18]:

- Random jump - generates a huge number of data points assuming uniform distribution of decision variables and selects the best one.
- Random walk - generates the trial solution with sequential improvements using scalar step length and unit random vector.
- Random walk with direction exploitation – is the improved version of random walk search, the successful direction of generating trial solution is found out and steps are taken along this direction.

5 Binary Character Recognition – Results

The learning and memory acquisition of optimization algorithms within OAT is verified through its application to a binary character recognition problem (Fig. 2). The aim is to examine the accuracy of binary character recognition for the analyzed algorithms: random search and parallel mutation hill climber.

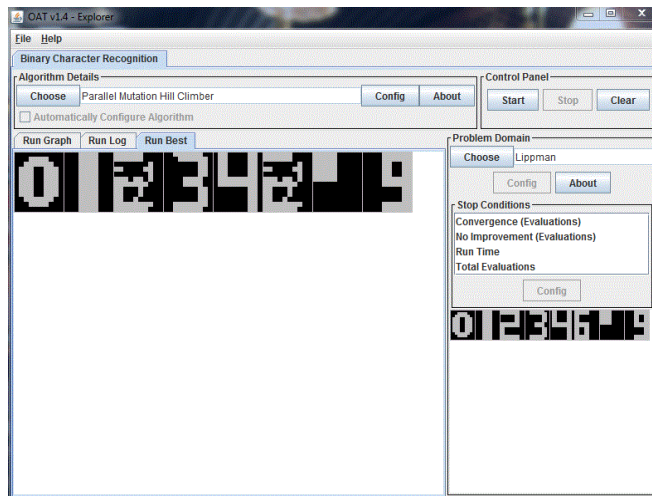


Figure 2

OAT environment for examination of binary character recognition

The parallel mutation hill climber algorithm (combined CLONALG – hill climber algorithm) - In the present instance, it can be supposed that the antigen population is signified by a set of eight binary characters to be learned. Each character is represented by a bitstring of length $L = 120$. The antibody configuration is

composed of 10 individuals, 8 of them in the memory set *M*. The initial parameters (set by OAT) are the following: mutate factor = 0.00833, population size = 16.

For different values of the population size (for the constant mutate factor), the output results (memory sets) are shown in Figure 3.

Similar to the previous case, with the choice of different values of the mutate factor, the resulting situation is shown in Figure 4. (the initial parameters being: mutate factor = 0.00833 (variable), population size = 30 (constant)).

Having in mind Figure 3 and Figure 4, two conclusions can be made. Through, increasing the number of population, with a fixed mutate factor, the quality of binary character recognition is also increased. In addition to this, the impact of the mutate factor (with a fixed population size) is not directly proportional. This practically means that it is necessary to examine several different values in order to obtain the best accuracy of recognition – with finding the optimal value.

Memory set	Population size
	Input patterns
	20 generations
	25 generations
	30 generations
	35 generations
	40 generations

Figure 3

Input and output patterns in function of the population size

Output	Mutate factor
	0.00833
	0.01
	0.02
	0.05

	0.055
	0.06
	0.07
	0.1

Figure 4

Output patterns in function of the mutation factor

After thoroughly analyzing the graphical results above, it can be concluded that the parallel mutation hill climber algorithm is very sensitive to the choice of parameters. This, in fact, means that the performance of this algorithm must definitely be checked with several different parameters, and based on the obtained, the best value can be chosen. This procedure is known in the literature as tuning parameters and directly affects the accuracy and the speed of convergence of the algorithm (e.g. solving time) to the solution. Also, it can be noticed that, due to the CLONALG component, this algorithm maintains a diverse set of local optimal solution, which can primarily be explained by the selection and reproduction schemes. Another important characteristic is the fact that CLONALG also considers the cell affinity, which corresponds to the fitness level of a given individual so as to determine the rate of mutation used for every member of the population.

In the case of a random search, the result is presented in Figure 5. For this type of algorithm, there are no configurable parameters and the quality of binary character recognition is very low.



Figure 5

Output patterns in the case of a random search

Conclusions

The practical part of the paper is focused on examining the impact of parameters of analyzed optimization algorithms on the accuracy of binary character recognition. In this sense, the existence of optimal value was found, as well as the extremely unacceptable values. This in fact means that great attention must be paid to the selection of parameter values (it is not to be disregarded which values are applied) and it is vital to test the behavior of the algorithm in practice with several different parameters. Moreover, it was confirmed that all of the optimization algorithms are not equally suitable for binary character recognition.

Namely, it has been shown that some of the analyzed algorithms did not generate an acceptable result.

References

- [1] The Nature of Mathematical Programming. Mathematical Programming Glossary. INFORMS Computing Society, at <http://glossary.computing.society.informs.org>
- [2] J. J. Weinman, E. Learned-Miller, A. R. Hanson: Scene Text Recognition Using Similarity and a Lexicon with Sparse Belief Propagation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*. 31(10), pp. 1733-1746, 2009, doi:10.1109/TPAMI.2009.38
- [3] M. Yokobayashi, T. Wakahara: Binarization and Recognition of Degraded Characters Using a Maximum Separability Axis in Color Space and GAT Correlation. in *Proc. of 18th Int. Conf. on Pattern Recognition (ICPR2006)* Vol. 2, 2006, pp. 885-888, doi:10.1109/ICPR.2006.326
- [4] T. De Campos, M. Barnard, K. Mikolajczyk, J. Kittler, F. Yan, W. Christmas, D. Windridge: An Evaluation of Bags-of-Words and Spatio-Temporal Shapes for Action Recognition. in *Proc. WACV. 2011*, pp. 344-351
- [5] Optimization Algorithm Toolkit, at <http://optalgtoolkit.sourceforge.net/>
- [6] R. Agrawal: A Modified Hill Climbing Algorithm. *International Journal of Emerging trends in Engineering and Development*. Issue 2, Vol. 6, 2012, pp. 663-667
- [7] Hill Climbing Methods, at <http://www.scribd.com/doc/7290255/Hill-Climbing-Methods>
- [8] W. Cohen, R. Greiner, D. Schuurmans, W. Cohen: Probabilistic Hill-Climbing. In S. Hanson, T. Petsche, R. Rivest, M. Kearns eds. *Computational Learning Theory and Natural Learning Systems*. Vol. II, MITCogNet, Boston: 1994
- [9] B. P. Gerkey, S. Thrun: Parallel Stochastic Hill-climbing with Small Teams. in L. E. Parker et al. (eds.) *Multi-Robot Systems: From Swarms to Intelligent Automata*. Volume III, Springer, 2005, 65-77
- [10] S. Forrest, M. Mitchell: Relative Building-Block Fitness and the Building-Block Hypothesis. In D. Whitley (ed.) *Foundations of Genetic Algorithms 2*, Morgan Kaufmann, San Mateo, CA, 1993
- [11] L. N. de Castro, J. Timmis: *Artificial Immune Systems: A New Computational Intelligence Approach*. Springer. 2002 Edition, pp. 57-58
- [12] L. N. de Castro, F. J. Von Zuben: Learning and Optimization using the Clonal Selection Principle. *IEEE Transactions on Evolutionary Computation*. 2002 Jun, 6(3):239-251

- [13] I. Aydin, M. Karakose, E. Akin: Chaotic-based Hybrid Negative Selection Algorithm and its Applications in Fault and Anomaly Detection. Expert Systems with Applications. Volume 37, Issue 7, pp. 5285-5294, 2010
- [14] COPASI Development Team, at <http://www.copasi.org>
- [15] L. A. Rastrigin: The Convergence of the Random Search Method in the Extremal Control of a many Parameter system. Automation and Remote Control. 24 (10), 1963, pp. 1337-1342
- [16] P. Cisar, S. Maravic Cisar, B. Markoski: Implementation of Immunological Algorithms in Solving Optimization Problems. Acta Polytechnica Hungarica. Vol. 11, No. 4, 2014
- [17] Hill Climbing, at <http://stackoverflow.com/tags/hill-climbing/info>
- [18] Advanced Topics in Optimization, Direct and Indirect Search Methods, at http://nptel.ac.in/courses/Webcourse-contents/IISc-BANG/OPTIMIZATION METHODS/pdf/Module_8/M8L4slides.pdf

Economic Aspects of Multi-Source Demand-Side Consumption Optimization in the Smart Home Concept

Aleš Horák¹, Miroslav Prýmek¹, Lukáš Prokop², Stanislav Mišák²

¹ Masaryk University, Botanická 554/68a, 602 00, Brno, Czech Republic

² VSB – TU Ostrava, 17. listopadu 15, 708 33, Ostrava, Czech Republic

hales@fi.muni.cz, mjp@mail.muni.cz, lukas.prokop@vsb.cz,

stanislav.misak@vsb.cz

Abstract: Current energy consumption trends lead to rapidly growing consumption of local renewable energy sources. Such installations bring new requirements on energy consumption profiles. Due to the massive multiplication of the results, one of the most interesting elements of the power grid, in this respect, is formed by households. Smart profiling of household energy consumption may be crucial for the adaptability of the global grid. In this article, we present the design and usage of a demand-side, consumption profiling system named the Priority-driven Appliance Control System (PAX). We describe the main features of the PAX system and show its application using real-world data. The main benefits are presented as direct economic assets in connection with various household energy sources (energy grid and photovoltaic panels) and efficient usage with regard to government energy grants.

Keywords: energy consumption; renewable energy sources; demand-side management; Priority-driven Appliance Control System; PAX

1 Introduction

A significant shift from traditional energy generation systems, to energy systems with distributed energy production and widespread integration of renewable energy sources in modern power systems, often called “smart grids”, present new challenges for research and development. Balancing power demand and supply in a distributed energy system, with a large number of installed types of renewable energy sources (RES), is an important aspect of a smart grid systems.

Power demand-side management (DSM) techniques and other balancing methods, under the smart grid environment, are generally successful with optimization issues at various energy levels [20]. In this paper, we focus our research on the

smallest of the smart grid unit, a “smart home”. The smart home concept as a smart grid elementary unit was chosen here to describe the principles, interconnections and relationships of DSM. The application of DSM techniques to the optimization of smart home appliance switching scheme according to local renewable energy production is one of the main goals of this paper.

The demand-side management techniques have been used for various purposes. The output of renewable energy generators (such as wind and photovoltaic) varies with weather conditions and generally it is not straightforward to modulate the output of RES to follow a particular load shape. An overview of current methods for the demand-side management can be found in [1], [18], [19]. The main DSM techniques include night-time heating with load switching, direct load control, load limiters, commercial/industrial programs, frequency regulation, time-of-use pricing or demand bidding [13], [14].

A DSM technique presented in [12] reshapes the demand profile by applying a heuristic optimization algorithm and a Genetic Algorithm Optimization. Evolutionary soft-computing methods are applied in [2]. Another AI technique, an Artificial Neural Networks (ANN), is used to implement a short-term load forecasting module integrated with the proposed DSM technique to estimate the load profile for a 24 hour period. Neural networks are applied in [9] for the demand-side management in a household, with an installed photovoltaic (PV) system. The control system is composed of two modules: a scheduler and a coordinator, both implemented via neural networks. The control system enhances the local energy performance, scheduling home appliances switching and maximizing the use of PV system.

In [7], the demand-side management for smart home is introduced in three different approximations: in a Round Robin scheduling scheme, in a Highest Power Next (HPN) priority scheme and in a Reciprocal Fair Management (RFM) approximation. Both the Round Robin and the HPN scheme incorporate sorting algorithms and thus, need high computational complexity. In, a combination of real-time pricing with the inclining block rate model is used by adopting a combined pricing model. The proposed power scheduling method reduces the electricity cost and the “Peak to average ratio.” Since these kinds of optimization problems are usually nonlinear, genetic algorithms are used. A game theory was used in [11] to formulate an energy scheduling game, where the players are the users and their strategies are the daily usage of their household appliances and loads.

An appliance commitment algorithm that schedules thermostatically controlled household loads based on price and consumption forecasts is introduced in [5]. The formulation of an appliance commitment problem is described using an electric water heater load. The thermal dynamics of heating and coasting of the water heater load is modeled by physical models; random hot water consumption is modeled using statistical methods. The work in [8], presents a load management

technique for the air conditioner loads in large apartment complexes, based on systematic aggregations and load factor controls of the air conditioner loads based on a queuing system model and the Markov birth and death process. The proposed technique can offer effective and convenient load management measures to power companies and large customers. Techniques based on the multi-agent approach are exploited in [4].

Finally, our approach, focuses on a DSM technique for optimizing the daily load profile of a smart home, where various types of home appliances are installed. The primarily aim of this work is to demonstrate that intelligent demand-side consumption control, is not just a theoretical concept, but of practical importance and has a significant positive financial impact.

Advantages of our approach, are in the application of Erlag, especially because of flexibility of this language and possibility for the use of inexpensive devices for DSM. The main focus of this paper is to present results of the application of PAX on Green premium policy. The daily load optimization is processed according to the actual production of the installed photovoltaic system in order to achieve a maximum financial surplus. The details of the computation of the economic aspects of the system management are explicated with a particular example connected with the government photovoltaic subsidy named Green Premium. Details for this economic model and are presented in Section 4.

2 Smart Household Architecture

In this research, we have focused on the lowest level of the smart grid structure - a household. The main aims of the control logic at this level is to control the operation of particular appliances, according to the actual power supply and its economic parameters, to minimize demand peaks, to make the overall consumption profile as fluent as possible (by coordination of particular appliance operation), to flexibly react to external effects (outage, brown-out, etc.) to minimize possible losses and reach given goals without significant impact on user experience.

In the following text, we describe the *Priority-driven Appliance Control System (PAX)*. The PAX system is designed for small and economic microcontrollers (A prototype hardware implementation uses Atmel microcontrollers in the price of USD 5, see [16] for details.) for a particular appliance control, yet are flexible enough to meet the above-stated criteria. The controlling core of the system can mix real and virtual appliances, the system can thus be used as a smart home simulator, as well as, a real, physical appliance controller. In the following sections, we describe how real world data are used as a basis for smart home modeling.

2.1 Data Sources and Measurement

Long-period measurements (more than one year) were performed to collect actual household data to test the PAX system functionality. Each household appliance in selected real households was monitored using power network analyzer MDS-U [6] during a long time period. MDS-U is a power analyzer which is able to measure voltage, current and power factors and calculate electrical quantities like active, reactive and apparent power for selected a time interval. In this case, a one minute time interval was used for data collection. Averaged power consumption curves were defined, based on long time power consumption data, for the 20 most common household appliances (refrigerator, cook top, wall oven, personal computer, washer, dishwasher, vacuum cleaner, etc.). During the long-term measurement, the switching scheme was evaluated for all monitored household appliances. Usual power consumption for any daytime can be defined based on averaged power consumption curves and the switching scheme for the most common household appliances. The power consumption curves of household appliances are the fundamental data for PAX testing.

Together with household appliances monitoring, selected renewable power sources were monitored during a one year time period. According to the actual trends in renewable energy sources utilization, photovoltaic (PV) and wind power plants (WPP) were chosen. PV consists of mono crystalline panels Aide Solar ($P_{MAX} = 180$ W, $I_{MP} = 5$ A, $U_{MP} = 36$ V, $I_{SC} = 5.2$ A, $U_{OS} = 45$ V, and 5 kWp rated power). WPP uses synchronous generator with permanent magnets of 12 kV·A, voltage 560 V, current 13.6 A, torque 780 N·m and 180 rpm. A detailed description of the small off-grid power system test bed can be found in [10] or [14]. Scenario II was applied for PV system with mono crystalline panels BENQ Green-Triplex PM245P00 ($P_{MAX} = 250$ W, $I_{MP} = 8.17$ A, $U_{MP} = 30.6$ V, $I_{SC} = 8.69$ A, $U_{OS} = 37.4$ V, and 5 kWp rated power).

In the current PAX experiment, two data sources were used - a data set of power consumption and a data set of power production. Both the mentioned data sets are based on the one year measurements of real data.

2.2 The PAX System Implementation

The PAX system is designed as a multi-agent system based on agent communication between the core scheduler and the source and consumer agents.

The actual system implementation is developed using the Erlang programming language [3], which is excellent in processing large amounts of discrete data. Erlang is a functional language with no shared memory. The system consists of many autonomous function units which communicate with each other, by message passing only. That is why, applications written in Erlang, can be parallelized in a

natural way and can run easily on multicore processors or even separate machines. The parallelization scalability is almost linear.

Erlang also has very good support for runtime code compilation and knowledge base changes which is used extensively for particular virtual appliance control in the PAX system and for the user interface adaptations.

2.3 Online and Offline Usage

As we have mentioned above, the PAX system is designed to support the development of smart home automation from the simulation phase up to device control in real-time. For the later usage scenario, the scheduler is driven by real-time clock, the events are induced by external sources or occur at predefined times.

When used as a simulator, without connections to real devices, there is no need to restrain the process with real-time clock. The (simulated) events should occur and be processed as fast as possible to make mass data processing possible. For this purpose, the PAX core implements an event queue. The queue is filled with simulated or real-world measured data and the simulation consists of sequential processing of the queue events.

Each simulation step has several phases:

1. Fetch subsequent event from the event queue
2. Deliver it to the appropriate agent
3. The agent receives the event and/or
 - (a) Changes its inner state
 - (b) Inserts a new event into the event queue
4. The simulation ends when the event queue is empty

Since the PAX core heavily depends on (asynchronous) message passing, it is possible to convert almost any code sequence into a sequence of events in the event queue. The same code can thus be used for offline data processing (simulation) and online device control. With this pure event-driven design, it is possible to implement arbitrary time precision simulation and simulate every possible aspect of the control system.

2.4 Appliance Categories

The grid appliances are classified into several types according to their control mechanism, user expectations and power consumption profiles.

2.4.1 Interactive Appliances

An Interactive Appliance (IA) is directly controlled by the user and the IA's reaction cannot be deferred. Also the power consumption is mostly constant and cannot be controlled.

A typical example of such an appliance is electric light, electric kettle, cook top, garage doors, gate, iron, lights, microwave oven, mixer, vacuum cleaner.

2.4.2 Intelligent Interactive Appliances

Intelligent Interactive Appliance (IIA) is a special case of the previous type. The main difference lies in the IIA's power control. The appliance can have more power consumption profiles that can be applied according to the current situation. The appliance can be also driven by special extra communication/control treatment (e.g. a server computer power cannot be cut off immediately, instead a control signal must be emitted and the computer will undergo the internal shut down process as soon as possible).

A typical example of an IIA is a computer.

2.4.3 Deferrable-Operation Appliances

A Deferrable-Operation Appliance (DOA) is also controlled by the user but an immediate operation is not necessary. When the user commands the appliance to operate, he/she only expects it to begin the work in a reasonable (configured) amount of time. The user does not depend on the precise time of the operation start and end, only the operation result must be delivered appropriately.

Typical examples of DOAs are a washing machine, dryer, dishwasher, slow cooker, car battery charger and generally all appliances which require some user intervention before operation and whose products are expected to be available for a longer time period.

2.4.4 Feedback-Controlled Appliance

A Feedback-Controlled Appliance (FCA) is usually designed to keep a predefined and (repeatedly) measured value within specified limits. The value is spontaneously tending in one direction and the power supply is needed to push it in the opposite direction. A conventional operation cycle is as follows: whenever a feedback value reaches the lower bound, the appliance engine is powered up to push it to the opposite bound and powered off as soon as this value is reached. This way the power profile of the appliance consists of alternating periods of maximum and none power consumption of a more or less constant duration.

A generalization of the FCA model in PAX is based on the assumption that the boundaries of the appliance operation should not be defined as hard limits that cannot be exceeded but rather as gradually increasing measures of the power demand priority (see the next section for details). From this point of view, a FCA can be seen, as a highly flexible device whose power consumption can be efficiently planned to fit together with a power consumption profile of other less-controllable devices. From the most general point of view it can be seen as a power storage device. Of course, FCA cannot store arbitrary amount of power - extreme values of stored power would have unintended consequences such as e.g. food going to rotten in an under-cooled refrigerator. A typical example of a FCA is a gas boiler, refrigerator, oven or steam cooker.

2.5 From Real Devices to Virtual and Back Categories

A development of a successful energy consumption control scheme consists in several steps:

1. Measure power consumption profiles of typical household appliances
2. Pre-process the data to remove any non-significant fluctuation (at a desired precision level)
3. Convert the data into power consumption change events
4. Test different consumption planning algorithms on the simulated appliances
5. Test the best algorithm on the real appliances

The main advantage of the PAX system is that the steps 2 to 5 can be done using the same software system thus keeping the simulation results as close to the real deployment as possible. Typical consumption profiles and parameters will be implemented into PAX in order to make the installation more user friendly, but several steps must be optimized for each installation.

3 Model Construction – Accumulator Type Appliance

In this section, we illustrate the above described process of transferring actual appliance measured data, into a virtual appliance model that is suitable for the PAX simulation.

Figure 1 displays parts of actual power consumption profiles of a typical boiler (central water heater) and a Television, measured with a one minute resolution.

The exact measured data are represented by the grey line. However, for model development, the data are transformed to a “cleaned” curve (the black line).

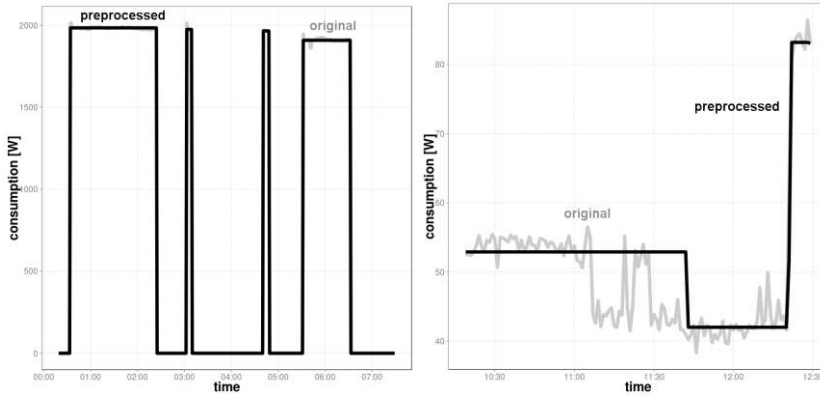


Figure 1

Measured power consumption and preprocessed profiles of a boiler (left) and a TV set (right).

The main objective of the data cleaning is to identify the consumption level breaking points which are then translated to consumption-change events with specific time values. Such events are included into the model communication flow, either in a form of a repeating pattern or (after statistical analysis) reproduced on a (partly) random basis.

The consumption data preprocessing algorithm is as follows:

1. Set a relative threshold value (typically, 1/5 of the maximum consumption)
2. Find the next time point where the difference between the last mean value and an actual consumption exceeds the threshold... This is the next breaking point
3. Compute the mean value between the last two breaking points
4. Repeat until end of data

The number of breaking points (and thus the similarity of the original and preprocessed curves) is determined by the chosen threshold, which is manually tuned to achieve the desired accuracy of the model.

The power consumption profile of a boiler consists of periodic repeating peaks of the same height (consumption) and a covariate length and distance. This behavior is caused by the thermostat, which keeps the inner temperature of the boiler within given limits. Generally, we consider the boiler to be an instance of a feedback-controlled device which acts like a perpetual energy accumulator. Whenever the accumulated energy falls below a given value, the device consumes the electrical power from the grid and accumulates it up to the upper limit value.

Table 1

Events generated from preprocessed measured data for the boiler model: $\{\{\{date\}, \{time\}\}, \{message: action, appliance, priority, consumption\}\}$

```

{{{2013,03,04}, {0,34,0}}, {set_consumption, boiler1, 4, 1.984}}
{{{2013,03,04}, {2,25,0}}, {set_consumption, boiler1, 0, 0}}
{{{2013,03,04}, {3,03,0}}, {set_consumption, boiler1, 4, 1.979}}
{{{2013,03,04}, {3,10,0}}, {set_consumption, boiler1, 0, 0}}
{{{2013,03,04}, {4,41,0}}, {set_consumption, boiler1, 4, 1.970}}
{{{2013,03,04}, {4,49,0}}, {set_consumption, boiler1, 0, 0}}

```

A conventional boiler does not communicate with its environment and follows the accumulation cycle forever in the safe range. Such behavior brings three main disadvantages:

1. If there are several appliances with a similar behavior, they produce unwanted random consumption peaks when their accumulation phases overlap
2. In case of an emergency state of the grid (blackout or brownout), such a device can only be switched off. There is no intermediate state which would not substantially impact the function of the device (to keep the water hot) while lowering the load on the power grid
3. The appliance does not allow, in advance, for the scenario of an anticipated/planned power loss

To fully exploit the potential of the smart grid concept, the PAX model of a smart boiler derived from real data is presented further. The smart boiler changes its behavior according to the state of the grid and cooperates with other power consuming devices.

In a virtual intelligent boiler model, based on real data, we must first identify the breaking points in the measured power profile of a real boiler and fit the profile to the fluctuation of the controlled value (the inner water temperature, i.e. the accumulated energy). The measured boiler water temperature oscillates between 49 °C and 58 °C. At 49 °C the heating system is switched on and works for 5-110 minutes until the temperature of 58 °C is reached (see Figure 1). The idle phase then lasts 30-90 minutes. Moreover, in the measured household, the boiler operation is limited by the time period of low-tariff energy lasting approximately from 20:30 till 6:30 (with flexible gaps summing to two hours in between). The heating speed can be approximated as 0.0017 °C per second. Of course, the real speed of heating is not linear but a linearization is a sufficient approximation in this case.

As mentioned above, the PAX system consumption scheduling is based on priorities of particular power consumption requests and sources (for details see [17]). The conventional nonintelligent boiler operates steadily in the given

temperature range and therefore its consumption requests have only one given priority P , where P is a high priority, e.g. $P = 4$ in our model. The appliance agent then periodically asks for power consumption of 1982 W with the priority 4 for 5-90 minutes (depending on the preceding idle period length). In the model input data, the measured values from Figure 1 correspond to the events presented in Table 1. For each time moment, the actual consumption of the appliance boiler1 is set to the value of ≈ 1982 watts with constant priority 4 or to the value of 0 watts (with 0 priority).

We have obtained a model of a standard boiler. With the inter-appliance communication ability brought by the PAX smart grid, we turn it into a flexible model of boiler behavior. The new model is based on the idea that with a decrease of the accumulated energy value (in the boiler case it equals the decrease of the inner temperature), the urgency of the power consumption increases because of a risk of boiler malfunction. In the PAX model this means that the priority of the consumption requests is inversely proportional to the boiler water temperature.

The temperature range is divided into priority zones, according to the risk that the device could miss its primary purpose (i.e. the risk of the water to cool down). If we approximate the temperature change function as stated above, the device agent can in each time interval estimate the time when the temperature crosses the priority border (if the actual consumption does not change). The priority ranges and the consumption prediction, is illustrated in Figure 2.

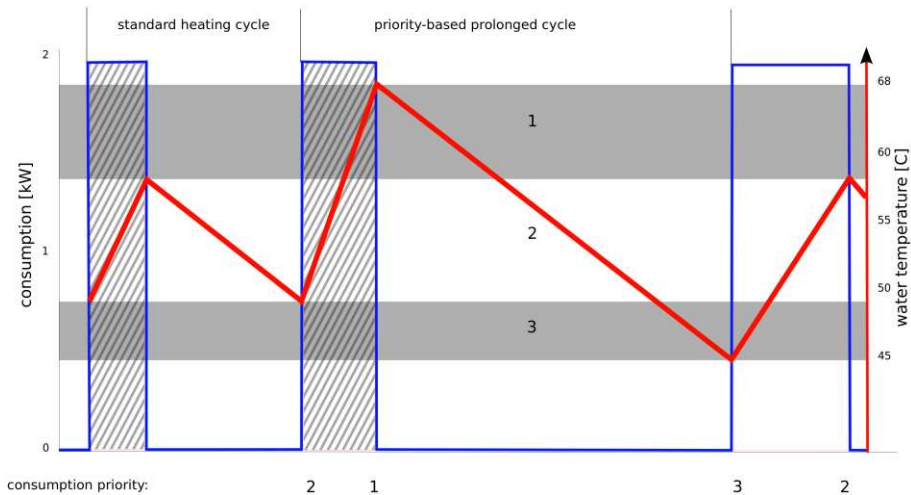


Figure 2
Priority ranges for boiler

At start, the agent requests the standard consumption time frame from the scheduler. If this request is denied, the agent computes the time when (with no power supplied) the temperature will cross the next lower priority limit and

requests the consumption from that time with a new (increased) priority. If the requested consumption frame is granted, the temperature crosses the priority boundary and the next request will have a decreased priority.

4 PAX Application – Green Premium Optimization

In recent years, the number of distributed power generators increases rapidly as small local power sources are being installed. The main representatives of newly installed power sources are photovoltaic (PV) systems with an installed capacity of around 5 kWp and positioned on the roof of the houses.

Such PV systems can be operated in various operation regimes according to national regulation rules. A frequent choice is using standard feed-in tariffs. The regime that we work with in this section exploits a government subsidy called Green Premium or Green Bonus.

The Green Premium policy supports the situation when the electrical energy produced by the household PV system is directly consumed by the house hold (or the household owner facilities). Only momentary surplus energy is sold in free energy market. Thus, to reach maximum financial profit and investment efficiency, as much produced energy as possible should be consumed by the PV system owner.

In this respect, the main issue of such household system is to optimize the energy consumption by shifting the load to the time when the PV system produces electrical energy. The PAX system is watching for PV system production and then control operation of each home appliances according to switching scheme and switching priorities. The main goal is to maximize energy consumption during PV system energy production using PAX system.

Generally, this leads to an optimization problem, where the load shift is optimized based on home appliances switching priorities and the actual installed PV system production. In the following text, we present the results of using the PAX system to solve such optimization tasks for a particular case of real household appliances and power sources. The households here are a common family houses with the floor area about 200 m² (scenario II 160 m²) with 1 child and 2 adults living in the house (scenario II 2 adults, 3 little children). Usual home appliances serve as standard power consumers and 5 kWp PV system is installed on the roof for the both households. The PV system operates in the Green Premium policy.

4.1 Photovoltaic System Operation

In the presented optimization task, the first possibility is to operate the household using standard feed-in tariffs. The feed-in tariffs for electric energy produced by PV systems are defined by national regulations. The current tariffs in Czech Republic (where the measurements are performed) are defined by Energy Regulatory Office (ERU) price decision number 4/2012 (See Table 2), which includes a detailed specification of the government financial support for renewable energy sources.

Table 2
Feed-in tariff and Green Premium for photovoltaic (PV) system

Row	PV system service start		Installed power [kWh]		Feed – in tariff	Green Premium
	from	to	from	to	[CZK/MWh]	[CZK/MWh]
513	Jul 1 2013	Dec 31 2013	0	5	2990	2440
514	Jul 1 2013	Dec 31 2013	5	30	2430	1880

The feed-in tariff for the analyzed households is 3 410 CZK per 1 MWh, while standard electricity tariff (of a major energy distributor) for residential sector is 5723.82 CZK per 1 MWh (Value added Tax (VAT) included). The difference between the feed-in and standard tariffs is fixed and does not leave space for local consumption optimizations.

The other possibility is the Green Premium policy that is based on the assumption that the more renewable energy is consumed in the house, the higher is the overall financial profit of such setup. When a PV system works within the Green Premium regime, two electric meters are wired into the system. The first electric meter measures only the energy produced by the PV system and the second electric meter (4Q electric meter) measure bi-directional power flow from and to the power distribution network. The energy cost savings when following the Green Premium policy are twofold. First, the Green Premium subsidy is computed from the overall energy produced by the PV system. The current Green Premium (in the second half of 2013) for the measured households PV installation is 2440 CZK per 1 MWh. The second part of the Green Premium savings consists in the fact that the PV energy can be used as a part of the household energy consumption and the rest of it (the surplus energy) is sold under specific feed-in tariff for surplus energy, which is currently 640 CZK per 1 MWh.

4.2 Discussion

With an installation of a photovoltaic power source, the household energy system repeatedly changes the states between producing surplus energy that is being sold to the grid and consuming extra energy from the grid. The profitability of both

these modes of operation is graphically presented in Figure 3. As we can see here, each kilowatt hour of PV production is subsidized with 2.44 CZK, thus each kilowatt hour of surplus PV production corresponds to 3.08 CZK (the premium plus the sold energy market price).

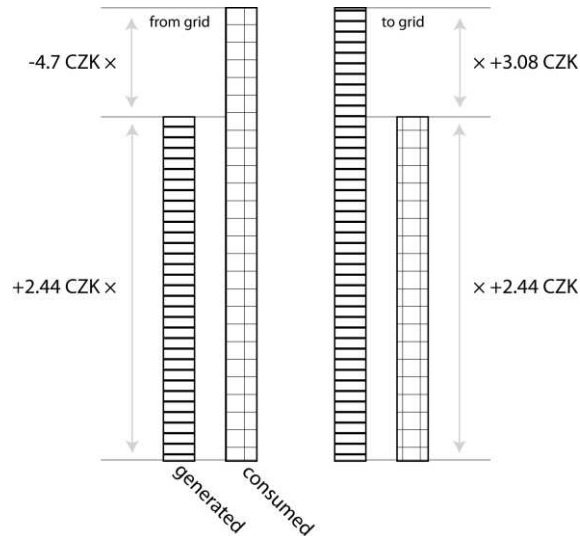


Figure 3

The make-up of the price of generated (from photovoltaic) and consumed (from grid) power. If a) the consumed power exceeds the generated power then the difference must be purchased for the standard price from the grid, otherwise b) the surplus power can be sold to the grid for the subsidized price.

In the time of lack of self-generated power, each kilowatt hour costs 4.70 CZK when purchased from the grid. The surplus energy can be thus processed in two possible ways: a) it can be sold for the prize of 0.64 CZK, or b) it can be consumed by an in house appliance that would otherwise consume expensive energy from the grid. This operation requires a time shift in the appliance operation. Each such kilowatt hour of consumption, when moved from the deficit to the surplus hour yields $4.70 - 3.08 = 1.62 \text{ CZK}$.

Therefore, the Green Premium policy offers a good opportunity for consumption profile optimization. The pressure for in-house consumption is here, even higher than the motivation for an energy savings (in the surplus hours).

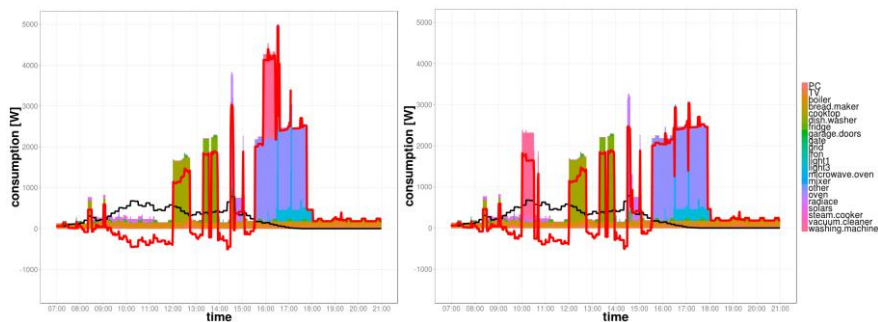


Figure 4

Comparison of the measured family house consumption profile with its optimized version

5 The Green Premium Model

On the side of the energy sources (photovoltaic and power grid), by setting the priority of the PV source higher than the priority of the grid power, the PAX system will automatically move the consumption into the area with unused PV power where possible. Of course the optimization improvement is highly dependent on the physical structure of the various appliances used within the household.

The DOA appliances category is the best category for the time-shifting optimization. Since we need to move as much consumption as possible to the PV output peak (noon), we must be able to defer the operation of the devices for several hours. This is possible only with DOAs hence this type of appliances yields the most of the achieved optimization.

From the Green Premium optimization point of view, the last two groups of appliances - IIA and FCA behave in a very similar manner. As we have seen in Section 3 and particularly in Figure 2, the operation of such appliances can be controlled intelligently to make a more suitable consumption profile but their operation generally cannot be postponed for hours. In spite of this, these appliances can be controlled with the aim of peak elimination. If there are many small areas of unused power distributed evenly throughout the PV generation peak (which surely is a case of household consumption profile), the PAX optimization reorganizes the overall consumption to fill the gaps. The resulting power consumption profile is more “flat”, with less peaks that need to be satisfied when power is used, from the grid.

5.1 Results and Discussion

The outcome of the optimization heavily depends on the inputs - especially on the number of the appliances in the household and their types. The type of the appliances varies significantly between particular households and substantially affects the consumption profile. The most influential, are the devices with high and steady power consumption such as e.g. electric heating. Such devices multiply the household power consumption and so there are special financial programs for households using these appliances. Therefore, the results presented in this paper cannot be generalized to such households.

Figure 4 shows the results of one day optimizations, of an analyzed household. The colored bars represent momentary consumption of particular appliances. The black line in the chart shows the daily course of the power from the PV system and the red line denotes the amount of energy either purchased from the global grid (> 0) or sold to the grid (< 0). We can see that most of a standard household consumption, does not follow the peak of the photovoltaic source. The aim of the optimization task thus lies in equalizing this discrepancy.

Table 3
Optimization Results

	Generated [kW·h]	Consumed [kW·h]	Purchased [kW·h]	Sold [kW·h]	Sold [CZK]	Overall Cash flow [CZK]	PV energy consumed	Cost reduction
Org. I	2865	5243	4027	1649	1055	10890	42%	0%
Opt. I	2865	5243	3770	1392	891	9847	51%	11%
Satur. I	2865	5243	2378	0	0	4190	100%	68%
Org. II	3156	5761	3952	1347	808	10074	57%	0%
Opt. II	3156	5761	3500	895	537	8220	72%	25%
Satur. II	3156	5761	2327	0	0	3241	100%	40%

Optimization of consumption by means of the PAX system was employed for a typical day with alternating cloudiness and a maximum supply of electric energy from PV around 20% of nominal power. These consisted of days during the autumn months where the element of diffuse radiation and PV making use of monocrystalline technology working with decreased effectiveness is predominant.

The much higher effectiveness of the PAX system can be recorded in the case of typical sunny days, in the spring and summer seasons. The direct element of sunshine, which PV with monocrystalline technology can make better use of, is predominant in these months. This positively manifests itself in a higher PV effectiveness. PV is used more often over the course of the daily cycle in the belt (70-100)% of the installed power and the length of the energy supply from PV is higher (approximately, 3 hours longer in comparison with a mostly cloudy day). As a result of the longer time interval for energy from PV and also a greater performance from PV, consumption can more effectively optimized with the use

of PAX, with the goal being, to purchase the least amount of energy, sell the least amount, in other words use as much energy from PV for one's own consumption.

Optimization by means of PV has higher effectiveness for this kind of day, which is given by the possibility of shifting, first and foremost, DOAs to the area of the maximum performance curve of PV. All of the appliances from DOAs are characterized by simple predictable cycles, and this including the stream profile and the length of duration. These can be fastened as the same time without actually significantly influencing the usual functioning of the household. These also consist of appliances which are energy demanding as they make up almost 70% of the appliances overall in the framework of the daily cycle. By means of appliances from the category DOAs, the profile under the power curve PV is thereby practically fulfilled at the time of its operations with a maximum effectiveness and a maximum performance. This fact results from the result of the amount of energy either purchased from the global grid (>0) or sold to the grid (<0), where there is an evident minimum amplitude within the period (12-16), in the afternoon.

Table 4

The electrical energy produced by the photovoltaic system (E_d means daily average during the month, E_m is overall power production in the month)

month	January	February	March	April	May	June
E_d [kW·h]	3.15	4.86	8.38	11.28	11.40	11.40
E_m [kW·h]	97	136	260	338	353	342
month	July	August	September	October	November	December
E_d [kW·h]	11.03	11.16	8.75	5.92	3.70	3.03
E_m [kW·h]	343	345	262	184	111	94

This is of course given by the fact that the majority of energy from PV is consumed whereby it is covered by the consumption of energy dominant appliances (dishwasher, washing machine, dryer). Apart from DOA appliances, the overall effectiveness of the energy system with the optimization of FCE can be increased, where the main representative in the analyzed case is a boiler.

For selected locality, the PV energy naturally varies not only during the day, but the variations also follow the course of seasons in the year. For the particular PV system the yearly history of average produced energy per month is shown in Table 4. Total energy production is 2865 kW·h.

The results of the PAX optimization of one-year consumption of the analyzed households are summed up in the Table 3. In the first row, there are original values taken from the measurement of the given household consumption profile. As we can see, the PV power is about one half of the household consumption but nearly 75% of the consumed power must be purchased from the grid operator at high price and more than a half of the PV production is sold for eight times lower market purchase price. The overall cost of the consumed power if it were

completely purchased from the global grid (i.e. without the PV installation), would be 24654 CZK (scenario I) and 27077 CZK (scenario II). The PV installation alone reduced this amount with the Green Premium subsidy of 6991 CZK (7701 CZK scenario II.), the 5718 CZK (6893 CZK scenario II) reduction in the energy purchased from the grid and the surplus PV power sold for the price of 1055 CZK (808 CZK scenario II). The overall cost with the PV installation thus drops to 10890 CZK (10074 CZK in scenario II).

These figures also show that in such typical household with the given PV size, the Green Premium subsidy makes the PV investment mostly, a partial coverage of the electrical power costs. This is consistent with the stress the subsidy puts on the local energy consumption, making it a good solution for individual family investment rather than for for-profit photovoltaic farms. From this point of view, the Green Premium program is surely a better solution over other kinds of PV subsidy programs.

The rows "Opt. I" and "Opt. II" of Table 3 shows the results of the model PAX optimization. As we can see, the amount of the purchased energy was lowered thus increasing the profitability of the PV installation. Time-shifting operations computed by PAX priority rules allowed an increase in the local consumption of the PV power to more than 51% (72% in scenario II). As described above, this number would be different for households with a different appliance structure, like Scenario II (more DOA appliances, such as a dryer used much more often because of care of 3 children).

Due to the PAX system ability to combine the real and virtual appliances, the financial impact of the appliances control system application is highly predictable. The predictability of the outcome is surely an important property for an overall PV investment/breakeven/profitability decision.

The results obtained by means of the PAX model reveal yet another possibility for increasing the PV investment profitability substantially under the Green Premium subsidy. The crucial figure here is the gap between the power purchase and sale prices. The part of the generated power that is left unused even after the consumption optimization represents a next big opportunity for the PV profitability improvement. This surplus power can be utilized by the last type of the appliance - the Intelligent Interactive Appliance (IIA). We can estimate the maximum financial effect of such an appliance in this way: let us suppose that we have a device which can be switched on in the time of the surplus power, has a scalable, power consumption and its activity has a positive financial effect. For the conservative estimation, we can suppose the financial effect is the same as the price of the consumed power - i.e. under standard circumstances; its action has zero profitability. But with the Green Premium subsidy, such a device will substantially improve the PV financial balance. Thanks to the PAX model, we can estimate the overlap between the power production and the consumption and thus estimate the remaining unused power to be sold to the network. The financial

effect of using all such energy by the described intelligent device is presented in Table 3, in the “Satur. I” row and Satur. II row. We can see that the financial potential of such energy usage is huge even with our conservative “normal zero profitability” assumption. Based on this, we can conclude that the development of such intelligent devices can be of high importance to Czech households using the Green Premium (or similar) subsidy programmer. Besides the above mentioned personal computer, a device of such kind can be represented by e.g. electric car charging stations, for-profit computation clusters or energy back-ups.

Conclusions

This paper introduced an application of the demand-side power consumption management principles that reshapes the demand profile by applying load shifting and load shedding, to reach a maximum financial benefit for the consumer.

The load shifting and load shedding operations are computed by the presented demand-side management system named PAX (*Priority-driven Appliance Control System*). The PAX system searches for runtime flexible optimal switching scheme for the household appliances according to appliance operation priority and energy sources availability.

We have presented, in detail, an application of the PAX system to the Smart Grid concept. As a particular example of direct exploitation of such a system, we have built a real-world model for optimizing the energy costs of a household equipped with a small on-roof photovoltaic system, operated under the Czech Republic government subsidy program called the Green Premium. This subsidy scheme is based on a principle that motivates the household owner to consume as much photovoltaic energy as possible, locally.

We have enumerated the energy costs, for an example household, in several analyzed scenarios. In the analyzed cases, the PAX system allowed savings of 9% of the photovoltaic power, without any negative effect on the user experience resulting in 11% for scenario I and 25% for scenario II. Application of the PAX system saved 1.043 CZK (11%) in scenario I and 1.854 CZK (25%) in scenario II. We have also discussed a possible setup directing photovoltaic power exploitations of up to a 68% (Scenario I) and 40% (Scenario II) reductions from the original costs.

In future research, we want to extend the application of PAX from Smart Homes to Smart Villages, where we suppose better utilization of the PAX advantages, due to more available appliances in *FCA* group and the related economies of scale.

Acknowledges

This paper was supported by the project LO1404: Sustainable development of ENET Centre, Students Grant Competition project reg. no. SP2015/170 and SP2015/178, project LE13011 and project TACR: TH01020426.

References

- [1] J. Aghaei, M. I. Alizadeh, Demand Response in Smart Electricity Grids Equipped with Renewable Energy Sources: A review, *Renewable and Sustainable Energy Reviews* 18 (2013) 64-72
- [2] R. Badawy, A. Yassine, A. Hessler, B. Hirsch, S. Albayrak, A Novel Multi-Agent System Utilizing Quantum-Inspired Evolution for Demand Side Management in the Future Smart Grid, *Integrated Computer-Aided Engineering* 20 (2) (2013) 127-141
- [3] F. Cesarini, S. Thompson, *Erlang Programming*, O'Reilly Media, 2009
- [4] L. Cristaldi, F. Ponci, M. Riva, M. Faifer, Multiagent Systems: an Example of Power System Dynamic Reconfiguration, *Integrated Computer-Aided Engineering* 17 (4) (2010) 359-372
- [5] P. Du, N. Lu, Appliance Commitment for Household Load Scheduling, *Smart Grid, IEEE Transactions on* 2 (2) (2011) 411-419
- [6] Egu Brno, s.r.o., Power Network Analyzer, available from http://www.egubrno.cz/sekce/s005/pristroje/mds/mds_ostatni_3_5_u.html
- [7] G. Koutitas, Control of Flexible Smart Devices in the Smart Grid, *Smart Grid, IEEE Transactions on* 3 (3) (2012) 1333-1343
- [8] S. C. Lee, S. J. Kim, S. H. Kim, Demand Side Management with Air Conditioner Loads Based on the Queuing System Model, *IEEE Transactions on Power Systems* 26 (2) (2011) 661-668
- [9] E. Matallanas, M. Castillo-Cagigal, A. Gutierrez, F. Monasterio-Huelin, E. Caamanno-Martín, D. Masa, J. Jimenez-Leube, Neural Network Controller for Active Demand-Side Management with PV Energy in the Residential Sector, *Applied Energy* 91 (1) (2012) 90-97
- [10] S. Misak, L. Prokop, Technical-Economical Analysis of Hybrid Off-grid Power System, in: 11th International Scientific conference Electric Power Engineering, 2010
- [11] A. H. Mohsenian-Rad, V. Wong, J. Jatskevich, R. Schober, A. Leon-Garcia, Autonomous Demand-Side Management Based on Game-Theoretic Energy Consumption Scheduling for the Future Smart Grid, *Smart Grid, IEEE Transactions on* 1 (3) (2010) 320- 331
- [12] M. Morgan, M. El Sobki Jr., Z. Osman, Matching Demand with Renewable Resources using Artificial Intelligence Techniques, in: *IEEE EuroCon 2013*, 2013
- [13] T. Pinto, I. Praca, Z. Vale, H. Morais, T. Sousa, Strategic Bidding in Electricity Markets: an Agent-based Simulator with Game Theory For Scenario Analysis, *Integrated Computer-Aided Engineering* 20 (4) (2013) 335-346

- [14] Pinto, T., Vale, Z., Sousa, T. M., Praça, I., Santos, G., Morais, H. (2014) “Adaptive Learning in Agents Behaviour: A Framework for Electricity Markets Simulation,” *Integrated Computer-Aided Engineering*, 21:4
- [15] Prokop, L., Mišák, S. Off - Grid Energy Concept of Family House, (2012) *Proceedings of the 13th International Scientific Conference Electric Power Engineering 2012, EPE 2012*, 2, pp. 753-758
- [16] M. Prymek, A. Horak, Modelling Optimal Household Power Consumption, in: *Proceedings of ElNet 2012 Workshop*, VSB Technical University of Ostrava, Ostrava, Czech Republic, 2012
- [17] M. Prymek, A. Horak, Priority-based Smart Household Power Control Model, in: *Electrical Power and Energy Conference 2012, IEEE Computer Society*, London, Ontario, Canada, 2012
- [18] P. Siano, Demand Response and Smart Grids - a Survey, *Renewable and Sustainable Energy Reviews* 30 (2014) 461-478
- [19] G. Strbac, Demand Side Management: Benefits and Challenges, *Energy Policy* 36 (12) (2008) 4419-4426, *Foresight Sustainable Energy Management and the Built Environment Project*
- [20] X. Wang, Y. Wang, Y. Cui, Energy and Locality Aware Load Balancing in Cloud Computing, *Integrated Computer-aided Engineering* 20 (4) (2013) 361-374

A New Model for Fine Turning Forces

Richárd Horváth

Óbuda University, Donát Bánki Faculty of Mechanical and Safety Engineering,
Népszínház u. 8, H-1081 Budapest, Hungary; horvath.richard@bgk.uni-obuda.hu

Abstract: Forces during turning depend not only on material properties and cutting parameters but to a great extent on the edge geometry of the tool as well, which determines chip shape (thickness and width). In fine turning it is almost exclusively the nose radius of the tool that does the cutting. The study reviews the main directions and results of researches in recent years concerning cutting force. It also presents the technology of fine cutting. Due to geometric considerations, chip characteristics are used that allow an exact description of cutting on nose radius as a function of the cutting parameters used. Dynamic tests are performed on two aluminium casting alloys and a mathematical model is constructed specifically for fine turning, using which expected forces can be estimated quite precisely during technological process planning.

Keywords: fine turning; force measuring; force model; regression; Kienzle model

Abbreviations

A – chip cross section, mm^2
 a_p – depth of cut, mm
 AS12 – die-cast eutectic aluminium alloy
 AS17 – die-cast hypereutectic aluminium alloy
 b – chip width, mm
 C, q, y , – constants of the regression equation
 ε_r – nose angle, °
 f – feed, mm
 F_c – main (tangential) force, N
 F_f – feed (axial) force, N
 F_p – passive (radial) force, N
 h – chip thickness, mm
 HB – Brinell hardness
 h_{eq} – equivalent chip thickness, mm
 HRC – Rockwell hardness
 HV – Vickers hardness
 k – specific cutting force, N/mm^2
 $k_{l,0.1}$ – constant value of main specific cutting force in new model, N/mm^2 (where $h_{eq}=0.1$ mm and $l_{eff}=1$ mm)
 $k_{l,1}$ – constant value of main specific cutting force in the Kienzle model, N/mm^2 (where $b=1$ mm and $h=1$ mm)
 l_{eff} – the cutting length of the edge of the tool, mm
 r_ε – nose radius, mm
 v_c – cutting speed, m/min
 κ_r – side cutting edge angle

1 Introduction

The examination of the dynamic behaviour of cutting processes can be performed with two methods: models can be deduced from mathematical calculations and empirical models. An early example of the former is the work of Merchant [1], whose dynamic model for orthogonal cutting is still part of research projects nowadays. In the past decades new materials, technologies and tools have appeared and so it has become an especially important research area. Many researchers have studied the machinability of various materials in recent years.

Suresh et al. [2] investigated the dynamic characteristics of AISI 4340 steel in turning operations with chemical vapor deposition (CVD) coated (TiC/TiCN/Al₂O₃) hard metal tool. Based on their results they obtained linear equations for the calculation of the resultant cutting force and the specific cutting force. They also concluded that cutting force and specific cutting force are most influenced by feed, less influenced by depth of cut, while they were least affected by cutting speed. Rao et al. [3] investigated the turning of AISI 1050 tempered steel – carbon steel (hardness: 484 HV) with a ceramic insert (Al₂O₃+TiC; - KY1615 -). They described the main effects with empirical formulae, performed a significance test of the parameters set for the cutting force, and made statements about the optimum. Szalóki et al. [4] investigated the cutting forces and surface roughness created in case of trochoidal milling. Their examinations were carried out using 40CrMnMo7 (1.2391) steel with a monolith cemented carbide HPC shank milling cutter. In their research they used design of experiment (DOE) and as a result an empirical equation depending on cutting parameters was set up to estimate cutting force components and maximum height of roughness profile R_z . Tállai et al. [5] tested thread formers with 5 different types of coatings on 40CrMnMo7. In their work they studied torque and wear of thread formers depending on formed length and compared the flood type and the minimal quantity lubrication (MQL). It is stated, that the increase of forming speed drastically decreases the number of machinable holes and wear condition of the forming tools can be followed by the torque measurement very well. The value of forming torque may be by 19 percent lower in case of MQL, compared to the flood type of lubrication.

In the past few years numerous researchers have studied the dynamic characteristics of hard turning. Aouici et al. [6], for example, turned AISI H11 steels of various hardness (40, 45 and 50 HRC) with a polycrystalline cubic boron nitride (CBN) tool. To calculate the individual force components, they used a quadratic equation, which contained the HRC hardness of the workpiece in addition to the usual data (v_c – cutting speed, f – feed, a_p – depth of cut). Their study also shows that cutting speed had the least effect on the force components. Lavlani et al. [7] also performed hard turning. In their study they turned MDN250 material (corresponding to 18Ni(250) maraging steel) with a coated ceramic tool (TNMA 160408S01525 - CC6050). They constructed a linear model for the three

measured force components, then they also stated that cutting force was mostly affected by feed and depth of cut, while cutting speed did not influence it significantly. Bouacha et al. [8] hard turned AISI 52100 bearing steel (64 HRC) with a CBN tool. They created second order phenomenological models for the estimation of the individual force components. Gaitonde et al. [9] studied the machinability of AISI D2 cold work tool steel. They performed their experiments with conventional and Wiper geometry ceramic tools (CC650, CC650WG and GC6050WH). They only considered depth of cut and time spent cutting in their investigation. They concluded that in the case of Wiper geometry tools cutting force increased linearly as depth of cut was increased. On the other hand, cutting force in the case of a conventional geometry tools increased up to a concrete value of depth of cut of $a_p=0.45$ mm, then started to decrease above this. Kundrák et al. [10] examined the microhardness of hard turned surfaces. They concluded that although according to the literature, cutting force does not have a direct effect on the hardness of the surface, indirectly it still has an effect because of the thermal energy, that the mechanical energy transfers into. Since the kinematics, geometry and technological parameters of hard turning often differ from those of conventional longitudinal turning, to determine cutting force with experimental formulae, further studies are necessary. Sztankovics et al. [11, 12], for example, presented how parameters characteristic of chip cross-section can be determined in the case of rotational turning.

In order to understand the behaviour and dynamic characteristics of the material during machining, in the case of the HSC (high speed cutting) technology, Pawade et al. [13] made experiments with softened Inconel 718 steel. They presented an analytical model, which predicts specific shear energy in the shear zone. They found that shear distances increase linearly as feed increases. Their model fitted their measurement results excellently.

In recent years more and more non-ferrous and non-metallic materials have been used in industry. Studying these has also got to the forefront. The Waldorf model [14, 15] of microcutting applies to the smallest chip thicknesses, because the undeformed (theoretical) chip thickness is less than 50 μm and is comparable to the edge radius of the tool. Annoni et al. [16] examined the machinability of C38500 (CuZn39Pb3) brass (hardness 81,5 HRB) with a hard metal tool (DCGX 070202–ALH10). In the range of microcutting they successfully changed the formulae for the calculation of cutting force and feed force because according to their results, the modified model better fits the values obtained during microcutting. Gaitonde et al. [17] also examined the machinability of copper alloy (CuZn39Pb3) (66 HRB). They performed their experiments with minimal quantity lubrication (MQL) with a hard metal tool of material K10 (TCGX 16 T3 08-Al H10). They varied cutting speed, feed and the amount of minimal lubrication (ml/h), while depth of cut was kept at a constant 2 mm value. They found that there is a considerable interrelationship between the amount of the lubricant and cutting speed. Machinability is very sensitive to the change of feed, irrespective of

the amount of the lubricant. They determined optimal cutting conditions, where specific cutting force (and average surface roughness, Ra) were minimal. Zebala and Kowalczyk [18] examined WC-Co material with a Mitsubishi triangular PCD tool (TNGA 160408). The cobalt content in the workpieces was 10, 15, and 25 wt%. Their research plan was based on the L_9 Taguchi method. Two empirical models were developed for the main cutting force (F_c). The first model was based on the power function; the second was based on the polynomial function according to modified RSM equations. When they compared the WC-Co with two different materials (Co contents 15 and 25 wt%), in terms of cutting force F_c , they found that there was no clear effect of Co content on the turning process. The lowest F_c values were obtained for the same cutting data. When they machined material with less Co content, higher F_c values for the same cutting data were generated. The analysis showed that depth of cut and cutting speed had the biggest influence on the F_c parameter.

Intensive research is also going on in the field of reinforced plastics and plastic-based composites. Hanafi et al. studied the turning of polyether ether ketone (PEEK) CF30 material with a TiN coated tool (WNMG 080408-TF) in dry conditions. They measured the three components of the force and determined the resultant force and specific cutting force with calculations. Empirical models were worked out depending on cutting data to predict both calculated forces. They tested the results with both response surface methodology (RSM) and Fuzzy algorithm, then compared the applicability of the two methods [19]. Fetecau and Stan [20] turned two kinds of polytetrafluoroethylene (PTFE) based composite materials: PTFE CG 32-3, containing 32% carbon and 3% graphite, and PTFE GR 15, containing 15% reformed graphite. They varied the cutting parameters (v_c , f , a_p) and used three polycrystalline diamond (PCD) tools of different nose radii. They found that feed and depth of cut had the greatest effect on cutting force, while the main component of cutting force was nearly constant as a function of cutting speed and the nose radius of the insert. An empirical equation was provided for both materials examined. The equations only contain the main parameters (feed, depth of cut).

The use of light metals is also becoming more and more widespread. De Agustina et al. [21] examined the dry turning of an aluminium alloy (UNS A97075). They used two tools of different nose radii (DCMT11T304-F2, DCMT11T308-F2) and measured the cutting force components, then compared the dynamic behaviour of the two tools. They found that at low feed the two tools of different radii required quite similar force. Joardar et al. [22] dry turned aluminium composites reinforced with SiC (LM6), using a polycrystalline diamond (PCD) tool. They constructed a second order model to estimate cutting force which also included Si content as input parameter in addition to cutting speed.

The author and his colleagues have already published articles on the machinability of die-cast aluminium parts, widely used in industry. They have examined in detail the surface roughness using different edge material and edge geometry.

They built phenomenological models to estimate surface roughness parameters and determined the optimum cutting conditions [23, 24, 25, 26]. They examined the changing of the statistical parameters of surface roughness separately as a function of cutting parameters, the machined materials, tool edge material and tool geometry [27]. This paper presents a new force model – describing the conditions of fine turning – by which, the main cutting force (F_c) can be calculated easily and accurately.

2 Materials and Methods

After iron and steel, aluminium and its alloys are used most widely in industry and they are only getting more and more popular. Due to their many advantages over pure aluminium, die-cast aluminium alloys (alloys of silicon, copper and magnesium) are widely used. Based on an experimental approach this paper analyses the dynamic conditions of fine turning of two generally used aluminium casting alloys.

2.1 Materials Used

Two aluminium casting alloys, used in great quantities in industry, have been selected. These alloys combine outstanding mechanical properties with technological advantages. The greatest advantage of the AS12 eutectic alloy is its excellent castability, while the AS17 hypereutectic alloy is harder than AS12 and more wear resistant.

The composition of the AS12 eutectic alloy is (wt %): Al=88.54%; Si=11.46%, while its hardness is 67 ± 2 HB_{2.5/62.5/30}.

The composition of the AS17 hypereutectic alloy is (wt %): Al=74.35%; Si=20.03%, Cu=4.57%; Fe=1.06%. Its hardness is 114 ± 3 HB_{2.5/62.5/30}.

The size of the workpiece used for the turning experiment was $\varnothing 110 \times 40$ mm.

2.2 Tools and Machine Tool Used

In their research so far the author and his colleagues have determined the optimum cutting conditions by inserting into their model not only the usual cutting parameters but the edge material and edge geometry as well [23, 24, 25]. The tool used for the dynamic test was the following: polycrystalline diamond (PCD), insert code: DCGW 11T304, edge geometry: ISO, manufacturer: TiroTool. The machine tool used was an EuroTurn 12B CNC lathe, with a spindle power of 7 kW and a maximum spindle RPM of 6000 1/min.

2.3 Cutting Force Measurement

Due to the conditions of chip removal and the low hardness of the materials, cutting forces are expected to be relatively small. The author developed and adapted a dynamometer system, the exploded view of which can be seen in Fig. 1. The sensitivity of the dynamometer was evaluated (pC/N), and its error curve was determined in the range from 0 N to 100 N. The dynamometer was connected to a KISTLER 5019 amplifier. Force data was evaluated with the DynoWare software. The cutting experiments were performed with a tool holder specially modified for measurement of small forces.

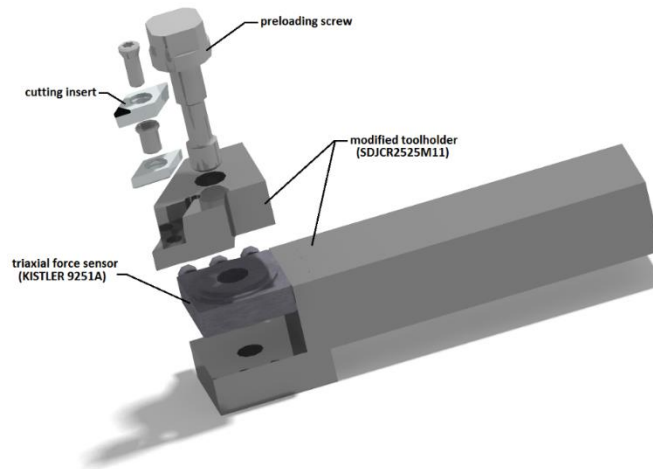


Figure 1
Exploded view of dynamometer system

2.4 The New Force Model Adapted to the Technology of Fine Turning

The Kienzle-Victor model [28] is still an excellent model widely used for the calculation of the components of cutting force. Specific cutting force (k) was introduced and determined by series of measurement, whose value depends on chip dimensions (theoretical undeformed chip width – b and chip thickness – h). In the case of a large chip cross-section (Figure 2), the side cutting angle of the tool (κ_r) provides a relationship between chip dimensions (b and h) and depth of cut (a_p) and feed (f) as cutting parameters. See Eqs. (1) and (2).

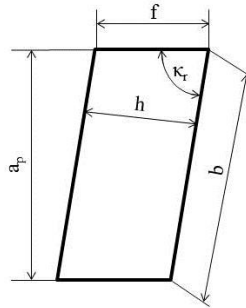


Figure 2

Chip cross-section for the Kienzle-Victor dynamic model

$$h = f \cdot \sin \kappa_r \quad (1)$$

$$b = \frac{a}{\sin \kappa_r} \quad (2)$$

The main value of specific cutting force ($k_{l,1}$) and the exponent of chip thickness (q) were determined experimentally. The main value of the specific cutting force mostly depends on the material and state of the workpiece and is true for chip dimensions $b=1$ mm and $h=1$ mm. The exponent also depends on the material but it is also influenced by cutting conditions. Thus the components of the cutting force can be calculated with the following general equation:

$$F = k_{l,1} \cdot h^{1-q} \cdot b \quad (3)$$

It has to be noted that the formula can only be used with correction coefficients in any case when real conditions differ from the experimental conditions (tool material, nose angle, wear, etc.)

The Kienzle-Victor method can be used well in the range of rough turning data, when depth of cut (a_p) is considerably larger than the nose radius of the tool (r_ϵ). Under fine turning conditions a smaller part of the side cutting edge and the whole of the nose radius takes part in chip removal. Therefore here chip geometry data h and b used by Kienzle and Victor are meaningless (Figure 2). It follows from this that characteristic $k_{l,1}$ cannot be used in the case of fine turning, either. That is why two novel chip characteristics were introduced (h_{eq} – equivalent chip thickness; l_{eff} – the cutting length of the edge of the tool) (Figure 3), which can describe chip geometry in fine turning accurately.

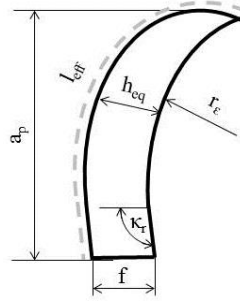


Figure 3

Characteristic chip cross-section in fine turning

The cutting parameters set, the side cutting edge angle (κ_r), and the nose radius of the tool (r_ϵ) mostly determine the cutting length of the edge of the tool (l_{eff}). Based on this concept chip cross-section can be given as follows:

$$A = a_p \cdot f = h \cdot b = h_{eq} \cdot l_{eff} \quad (4)$$

In fine turning effective edge length can be calculated with the formula [29, 30], which depends on a_p, f, κ_r (is given in radian), r_ϵ :

$$l_{eff} = \frac{a_p - r_\epsilon \cdot (1 - \cos \kappa_r)}{\sin \kappa_r} + r_\epsilon \cdot \left(\kappa_r + \arcsin \frac{f}{2 \cdot r_\epsilon} \right) \quad (5)$$

The equivalent chip thickness can be calculated from cutting parameters and effective edge length with the following formula:

$$h_{eq} = \frac{a_p \cdot f}{l_{eff}} \quad (6)$$

In fine turning $h_{eq} \ll 1$ mm is always true so $k_{1,1}$ cannot be used. Therefore it is worth introducing a method of calculation that describes chip geometry better. This is the main value of the specific cutting force in fine turning, $k_{1,0,1}$. It refers to $l_{eff} = 1$ mm and $h_{eq} = 0.1$ mm.

The newly introduced cutting force model requires the calculation of specific cutting force (determined with a dynamometer), which can be written as follows:

$$k = \frac{F}{A} = \frac{F}{h_{eq} \cdot l_{eff}} \quad (7)$$

The obtained k values depend on h_{eq} and l_{eff} ; therefore it is worth modelling them with a two-factor power function regression as below:

$$k = C \cdot h_{eq}^q \cdot l_{eff}^y \quad (8)$$

In Eq. (8) the fitting parameter C is a positive number, q and y are an empirically estimated exponents.

If the substitution $h_{eq}=0.1 \text{ mm}$ is used, the value of $k_{l,0.1}$ is the following:

$$k_{l,0.1} = C \cdot 0.1^q \quad (9)$$

And the resulting general cutting force model is:

$$F = k \cdot h_{eq} \cdot l_{eff} = k_{l,0.1} \cdot 10^q \cdot h_{eq}^{1+q} \cdot l_{eff}^{1+y} \quad (10)$$

2.5 Design of Experiments

The experiments used in this paper embraces the technological spectrum of fine turning ($f=0.03\text{--}0.15 \text{ mm}$; $a_p=0.25\text{--}0.7 \text{ mm}$). In the experiments cutting was performed at depth of cut $a_p=0.25 \text{ mm}$ where $a_p < r_\epsilon$, while at the highest value a small section of the side cutting edge also takes part in chip removal.

Researchers so far have reported that cutting speed has a negligible effect on cutting force [6,7,20], therefore speed is kept high as in industrial conditions but constant ($v_c=1000 \text{ m/min}$).

In the experiments it is reasonable to calculate the chip cross-section where $l_{eff}=1 \text{ mm}$ and $h_{eq}=0.1 \text{ mm}$, and determine the f and a_p values to be set in order to achieve it. Using equations (5) and (6):

$$l_{eff} = 1(\text{mm}) = \frac{a_p - 0.4(\text{mm}) \cdot \left(1 - \cos \frac{31 \cdot \pi}{60}\right)}{\sin \frac{31 \cdot \pi}{60}} + 0.4(\text{mm}) \cdot \left(\frac{31 \cdot \pi}{60} + \arcsin \frac{f}{2 \cdot 0.4(\text{mm})}\right) \quad (11)$$

$$h_{eq} = 0.1(\text{mm}) = \frac{a_p \cdot f}{l(\text{mm})} \quad (12)$$

Equations (10) and (11) yield the following pairs of solutions:

$a_p=0.699415 \text{ mm}$; $f=0.142976 \text{ mm}$ and $a_p=0.125135 \text{ mm}$; $f=0.799133 \text{ mm}$. It is easy to see that this requirement is fulfilled for $f=0.143 \text{ mm}$ and $a_p=0.7 \text{ mm}$. Table 1 shows the experimental settings for both materials. The set of experiments was made so that measurement point 22 is to determine (check) $k_{l,0.1}$.

Table 1
Experiment points

Measurement point	$a_p, \text{ mm}$	$f, \text{ mm}$	$l_{eff}, \text{ mm}$	$h_{eq}, \text{ mm}$	$A, \text{ mm}^2$
1.	0.25	0.03	0.493	0.015	0.0075
2.	0.25	0.05	0.503	0.025	0.0125
3.	0.25	0.07	0.513	0.034	0.0175
4.	0.25	0.09	0.523	0.043	0.0225
5.	0.25	0.11	0.533	0.052	0.0275

Measurement point	a_p, mm	f, mm	l_{eff}, mm	h_{eq}, mm	A, mm^2
6.	0.25	0.13	0.543	0.060	0.0325
7.	0.25	0.15	0.554	0.068	0.0375
8.	0.5	0.03	0.743	0.020	0.015
9.	0.5	0.05	0.753	0.033	0.025
10.	0.5	0.07	0.763	0.046	0.035
11.	0.5	0.09	0.774	0.058	0.045
12.	0.5	0.11	0.784	0.070	0.055
13.	0.5	0.13	0.794	0.082	0.065
14.	0.5	0.15	0.804	0.093	0.075
15.	0.7	0.03	0.944	0.022	0.021
16.	0.7	0.05	0.954	0.037	0.035
17.	0.7	0.07	0.964	0.051	0.049
18.	0.7	0.09	0.974	0.065	0.063
19.	0.7	0.11	0.984	0.078	0.077
20.	0.7	0.13	0.994	0.092	0.091
($k_{1,0,1}$) 21.	0.7	0.143	1.001	0.100	0.1001
22.	0.7	0.15	1.004	0.105	0.105

3 Results

3.1 Results of Main Force (F_c) Measurements

The experiments were repeated twice (the arrangement of force measurement can be seen in Figure 4).

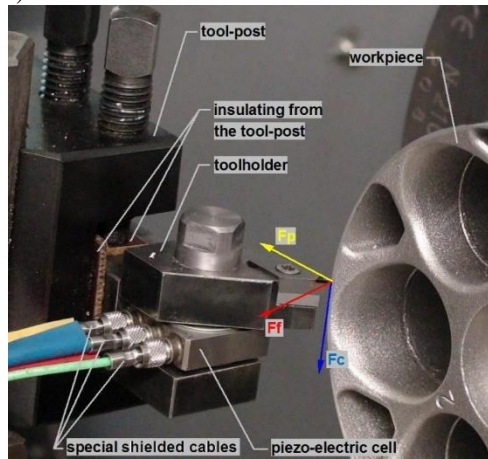


Figure 4

The layout of force measurement

Figure 5 shows specific cutting force as a function of equivalent chip thickness. It can be seen that in fine turning, too, the values of measurement points are on a straight line if the diagram is logarithmically scaled.

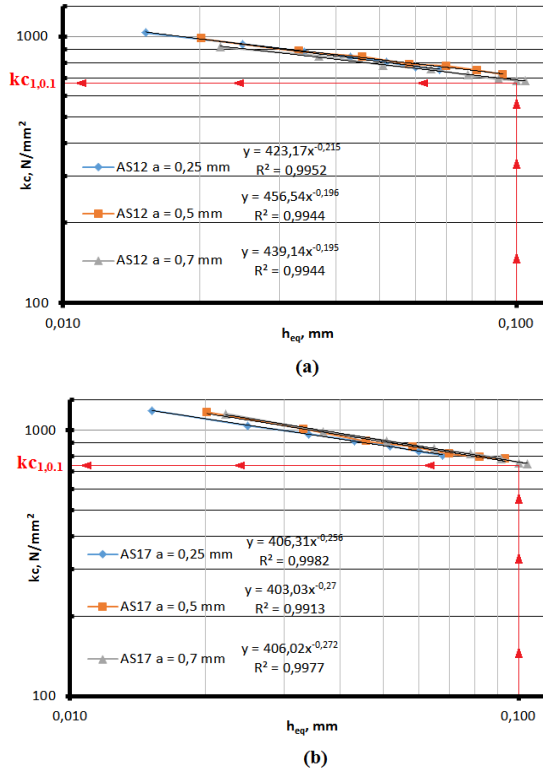


Figure 5

Specific cutting force as a function of equivalent chip thickness
a) in the case of AS12; b) in the case of AS17

Table 2 shows the averages of these F_c values, and the values predicted by the model (chapter 3.2.) and their relative deviation in %.

Table 2
Main cutting force values of AS12 and AS17 materials

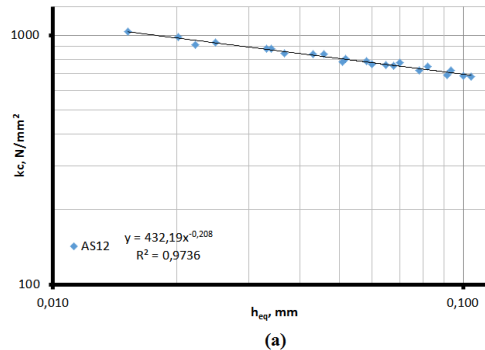
Measurement point	AS12				AS17			
	F_c measured <i>N</i>	k_c <i>N/mm²</i>	F_c calculated <i>N</i>	Error, %	F_c measured <i>N</i>	k_c <i>N/mm²</i>	F_c calculated <i>N</i>	Error, %
1.	7.77	1036.66	7.82	0.53	8.89	1185.26	8.97	0.91
2.	11.73	938.24	11.80	0.61	12.99	1039.00	13.11	0.92
3.	15.40	880.17	15.49	0.59	16.94	968.23	16.86	-0.47
4.	18.95	842.36	19.00	0.26	20.54	912.77	20.39	-0.71
5.	22.18	806.66	22.38	0.89	24.05	874.49	23.76	-1.19

	AS12				AS17			
Measurement point	F_c measured N	k_c N/mm^2	F_c calculated N	Error, %	F_c measured N	k_c N/mm^2	F_c calculated N	Error, %
6.	24.92	766.70	25.66	2.97	27.01	831.01	27.02	0.04
7.	28.23	752.75	28.86	2.22	30.17	804.56	30.19	0.05
8.	14.74	982.95	14.54	-1.40	17.55	1169.84	17.23	-1.80
9.	22.03	881.23	21.92	-0.48	25.22	1008.92	25.12	-0.42
10.	29.43	840.96	28.76	-2.30	31.79	908.23	32.24	1.42
11.	35.50	788.98	35.24	-0.76	38.99	866.39	38.90	-0.23
12.	42.65	775.43	41.46	-2.79	44.76	813.82	45.23	1.04
13.	48.58	747.39	47.48	-2.26	51.67	794.99	51.31	-0.70
14.	54.15	721.97	53.35	-1.47	58.45	779.39	57.21	-2.13
15.	19.22	915.10	19.77	2.90	24.15	1150.01	24.00	-0.64
16.	29.63	846.63	29.81	0.59	34.54	986.81	34.94	1.15
17.	38.34	782.41	39.08	1.93	44.73	912.96	44.80	0.15
18.	47.76	758.06	47.86	0.22	53.53	849.68	54.00	0.88
19.	55.46	720.22	56.29	1.50	63.04	818.66	62.73	-0.49
20.	63.14	693.85	64.44	2.06	71.29	783.39	71.10	-0.26
($k_{cI,0.1}$) 21.	68.78	687.07	69.62	1.23	75.52	754.43	76.39	1.16
22.	71.79	683.74	72.38	0.82	78.74	749.88	79.20	0.58

It can be seen from the Table 2 that the cutting force requirement of the AS17 aluminium alloy is greater than that of the AS12 alloy. This can be attributed to the differences in hardness, and the hard primary silicon grains forming in the AS17 hypereutectic alloy.

3.2 Model for Main Force (F_c)

Figure 6 a) and b) displays the specific cutting forces in a logarithmically scaled diagram as a function of equivalent chip thickness. It can be seen that the set of all the measurement points fit a straight line quite well. In Figure 6 c) and d) all measurements points are handled together. But specific cutting force is studied as a function of chip cross-section. It can be seen that the scatter of the values of k_c is great for both materials. It can be stated that specific cutting force depends far more significantly on the h_{eq} parameter introduced in the force model. This also explains the importance of h_{eq} .



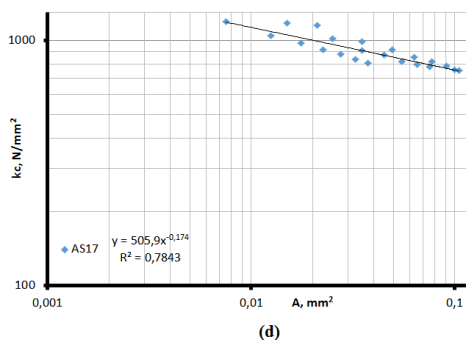
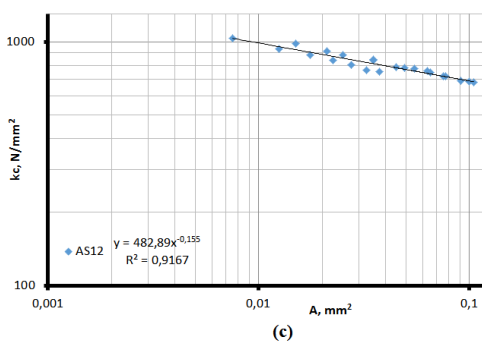
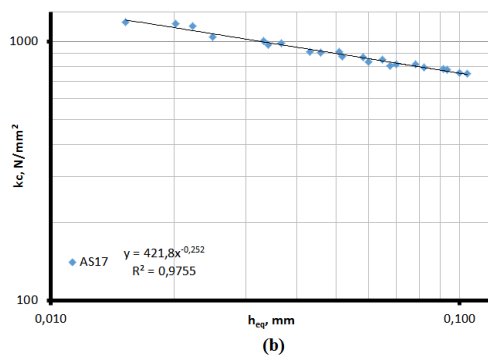


Figure 6

Specific cutting force as a function of equivalent chip thickness and chip cross-section

According to eq. (8) the equations of specific cutting force (depending on equivalent chip thickness and the cutting length) for the two materials are the following:

$$k_{c_AS12} = 439 \cdot h_{eq}^{-0,198} \cdot l_{eff}^{-0,039} \quad (13)$$

$$k_{c_AS17} = 408 \cdot h_{eq}^{-0,272} \cdot l_{eff}^{-0,088} \quad (14)$$

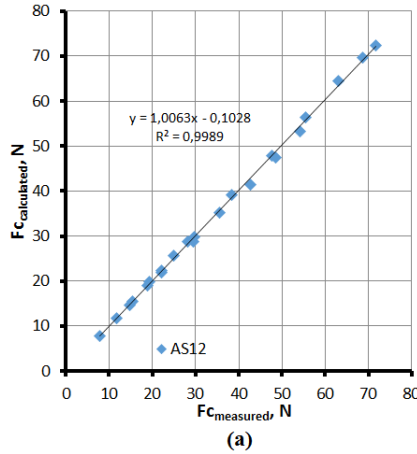
Equations (13, 14) yield that $k_{c1,0.1_AS12}=692 \text{ N/mm}^2$ and $k_{c1,0.1_AS17}=762 \text{ N/mm}^2$. After transformations based on Eqs. (7-10) we get the following equations for cutting force [31]:

$$F_{c_AS12} = 692 \cdot 10^{-0,198} \cdot h_{eq}^{0,8} \cdot l_{eff}^{0,961} \quad (15)$$

$$F_{c_AS17} = 762 \cdot 10^{-0,272} \cdot h_{eq}^{0,728} \cdot l_{eff}^{1,089} \quad (16)$$

3.3 Checking the Equations of F_c

The equations were checked by plotting and comparing the calculated and estimated cutting force values (Table 2) against the measured values. The deviation of calculated values from measured values in the case of AS12 is between -2.79% and 2.97%. In the case of AS17, deviation is between -2.13% and 1.42 %. The deviation of the plotted points from the $x = y$ line shows the good usability of the equations. The Figure 7 shows for both materials that the equations (eqs. 15, 16) approximate the measured values quite well.



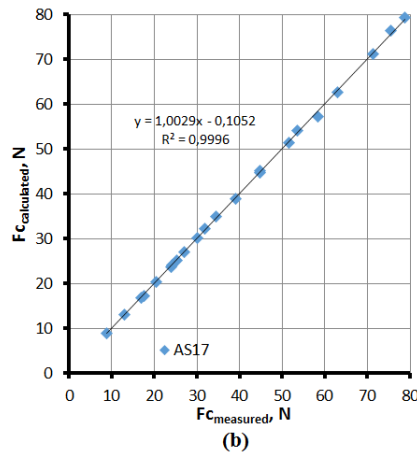


Figure 7

A comparison of measured and calculated roughness values in the case of (a) AS12 and (b) AS17

3.4 Models for the Feed Force (F_f) and Radial Force (F_p)

Table 3 and 4 show the averages of measured F_f and F_p values and the values predicted by the model (Table 5) and their relative deviation in%.

Table 3
Feed force values of AS12 and AS17 materials

Measurement point	AS12				AS17			
	$F_{f\text{ measured}}$ N	k_f N/mm ²	$F_{f\text{ calculated}}$ N	Error, %	$F_{f\text{ measured}}$ N	k_f N/mm ²	$F_{f\text{ calculated}}$ N	Error, %
1.	4.92	656.14	4.80	-2.53	6.11	814.87	6.05	-0.96
2.	5.92	473.86	5.83	-1.49	7.18	574.47	7.19	0.18
3.	6.56	374.97	6.63	1.01	8.03	459.02	8.05	0.27
4.	7.22	321.08	7.28	0.80	8.78	390.11	8.76	-0.21
5.	7.38	268.18	7.84	6.36	8.91	324.06	9.36	5.05
6.	8.12	249.98	8.34	2.63	9.64	296.53	9.89	2.63
7.	8.55	228.11	8.78	2.66	9.98	266.26	10.36	3.81
8.	5.61	374.18	5.71	1.67	7.24	482.91	7.30	0.80
9.	6.50	259.95	6.95	6.99	8.44	337.51	8.69	2.95
10.	7.62	217.69	7.91	3.83	9.68	276.70	9.73	0.50
11.	8.63	191.80	8.70	0.85	11.06	245.70	10.59	-4.20
12.	9.73	176.98	9.39	-3.54	11.23	204.17	11.33	0.89
13.	10.55	162.28	10.00	-5.24	12.34	189.91	11.98	-2.96
14.	11.76	156.77	10.54	-10.34	13.44	179.17	12.56	-6.52
15.	6.00	285.88	6.15	2.45	7.72	367.72	7.96	3.08
16.	7.22	206.32	7.50	3.84	9.23	263.71	9.47	2.63

Measurement point	AS12				AS17			
	$F_{f\text{ measured}}$ N	k_f N/mm ²	$F_{f\text{ calculated}}$ N	Error, %	$F_{f\text{ measured}}$ N	k_f N/mm ²	$F_{f\text{ calculated}}$ N	Error, %
17.	8.61	175.63	8.54	-0.80	10.53	214.97	10.62	0.80
18.	9.91	157.23	9.40	-5.10	12.02	190.75	11.56	-3.81
19.	10.92	141.85	10.15	-7.10	12.52	162.63	12.37	-1.24
20.	11.82	129.89	10.81	-8.56	13.66	150.16	13.08	-4.28
($k_{f1,0.1}$) 21.	12.97	129.52	11.20	-13.59	14.74	147.22	13.50	-8.36
22.	13.46	128.16	11.41	-15.24	14.79	140.90	13.72	-7.25

Table 4
Radial cutting force values of AS12 and AS17 materials

Measurement point	AS12				AS17			
	$F_{p\text{ measured}}$ N	k_p N/mm ²	$F_{p\text{ calculated}}$ N	Error, %	$F_{p\text{ measured}}$ N	k_p N/mm ²	$F_{p\text{ calculated}}$ N	Error, %
1.	4.01	534.44	3.87	-3.35	5.05	673.58	5.12	1.30
2.	4.88	390.77	4.71	-3.56	6.07	486.00	5.95	-2.06
3.	5.43	310.02	5.40	-0.53	6.73	384.66	6.62	-1.63
4.	5.65	251.29	6.00	6.18	7.39	328.22	7.21	-2.34
5.	5.76	209.52	6.56	13.90	6.83	248.37	7.75	13.53
6.	6.04	185.96	7.09	17.31	7.30	224.71	8.27	13.18
7.	6.41	170.83	7.60	18.56	7.49	199.86	8.76	16.82
8.	7.94	529.18	7.67	-3.36	10.50	700.13	9.91	-5.59
9.	9.39	375.68	9.26	-1.40	12.04	481.48	11.43	-5.00
10.	10.80	308.43	10.53	-2.42	13.50	385.83	12.63	-6.49
11.	12.07	268.14	11.64	-3.55	14.44	320.96	13.65	-5.49
12.	13.08	237.88	12.64	-3.41	14.61	265.72	14.57	-0.30
13.	14.33	220.47	13.56	-5.35	15.72	241.86	15.42	-1.91
14.	15.51	206.74	14.44	-6.88	16.64	221.81	16.22	-2.49
15.	10.42	496.14	11.16	7.07	13.39	637.69	14.32	6.95
16.	12.97	370.46	13.42	3.54	15.94	455.46	16.46	3.27
17.	15.15	309.14	15.22	0.50	18.27	372.79	18.12	-0.80
18.	17.54	278.36	16.77	-4.37	20.33	322.67	19.53	-3.95
19.	18.42	239.20	18.16	-1.42	20.55	266.93	20.77	1.08
20.	20.01	219.84	19.43	-2.85	21.86	240.24	21.92	0.26
($k_{p1,0.1}$) 21.	21.37	213.50	20.22	-5.39	22.76	227.35	22.62	-0.60
22.	22.19	211.38	20.63	-7.05	23.32	222.06	22.99	-1.41

The F_f and F_p force components can be estimated similarly to F_c by the new force model, which was introduced in chapter 2.4. The deviation of calculated values from measured values in the case of F_f :

- AS12 material: between -15.24% and 6.99%. In the case of AS17, deviation is between -8.36% and 3.81%.

In the case of F_p :

- AS12 material: between -7.05% and 18.56%. In the case of AS17, deviation is between -6.49% and 16.82%.

Although the deviations of F_f and F_p (force) components are higher than in case of F_c , those accuracies meet the requirements of technological process planning. The constants of the general equation (10) for calculating F_f and F_p are shown in Table 5.

Table 5
Constants of equations for F_f and F_p

	constants of the equation of passive (radial) force, F_p , N			constants of the equation of feed force, F_f , N		
	$k_{l,0.1}$	q	γ	$k_{l,0.1}$	q	γ
AS12	112	-0.393	0.153	202	0.340	1.430
AS17	135	-0.343	0.221	226	0.248	1.440

Conclusion

In this paper the forces of cutting on a widely used eutectic and a hyper-eutectic aluminium alloy with a diamond tool were examined. A new force model was introduced for the technology of fine turning, by which the main values of cutting forces can be estimated with ease and great accuracy. The following conclusions can be drawn about the new model and the test results:

- the model does not differ significantly from the widely used Kienzle-Victor formula, but operates with chip characteristic of fine turning;
- the specific cutting force formula (8) estimates specific cutting force well for both materials examined;
- the main value of the specific cutting force introduced in the novel model in case of fine turning ($k_{l,0.1}$, N/mm², where $l_{eff} = 1$ mm and $h_{eq} = 0.1$ mm);
- the main values of the specific cutting forces for the fine turning of aluminium alloys can be calculated easily with formula (9). (In the case of AS12: $k_{c,l,0.1} = 692$ N/mm²; $k_{f,l,0.1} = 112$ N/mm²; $k_{p,l,0.1} = 202$ N/mm². In the case of AS17: $k_{c,l,0.1} = 762$ N/mm²; $k_{f,l,0.1} = 135$ N/mm²; $k_{p,l,0.1} = 226$ N/mm²);
- it was verified that specific cutting force mainly depends on the equivalent chip thickness;
- From the test results it can be concluded that the equation of cutting force (10) predicts the measured force components (F_c , F_f , F_p) with great accuracy.

Acknowledgement

The author wishes to thank his colleague, Sándor Sipos, for his valuable advices and comments.

References

- [1] Merchant ME.: Basic Mechanics of the Metal Cutting Process, *Journal of Applied Mechanics* (1944) 68-75
- [2] R. Suresh, S. Basavarajappa, G. L. Samuel: Some Studies on Hard Turning of AISI 4340 Steel Using Multilayer Coated Carbide Tool, *Measurement* 45 (2012) 1872-1884
- [3] C. J. Rao , D. Nageswara Rao, P. Srihari: Influence of Cutting Parameters on Cutting Force and Surface Finish in Turning Operation, *Procedia Engineering* 64 (2013) 1405-1415
- [4] I. Szalóki, S. Csuka, S. Sipos: "New Test Results in Cycloid-Forming Trochoidal Milling." *Acta Polytechnica Hungarica* 11:(2) (2014) 215-228
- [5] P. Tállai, S. Csuka, S. Sipos: Thread Forming Tools with Optimised Coatings, *Acta Polytechnica Hungarica* 12:(1) (2015) 55-66
- [6] H. Aouici, M. A. Yallese, K. Chaoui, T. Mabrouki, J.-F. Rigal: Analysis of Surface Roughness and Cutting Force Components in Hard Turning with CBN Tool: Prediction Model and Cutting Conditions Optimization, *Measurement* 45 (2012) 344-353
- [7] D. I. Lalwani, N. K. Mehta, P. K. Jain: Experimental Investigations of Cutting Parameters Influence on Cutting Forces and Surface Roughness in Finish Hard Turning of MDN250 Steel, *Journal of materials processing technology* 206 (2008) 167-179
- [8] K. Bouacha, M. A. Yallese, T. Mabrouki, J.-F. Rigal: Statistical Analysis of Surface Roughness and Cutting Forces using Response Surface Methodology in Hard Turning of AISI 52100 Bearing Steel with CBN Tool, *International Journal of Refractory Metals & Hard Materials* 28 (2010) 349-361
- [9] V. N. Gaitonde, S. R. Karnik, Luis Figueira, J. Paulo Davim: Machinability Investigations in Hard Turning of AISI D2 Cold Work Tool Steel with Conventional and Wiper Ceramic Inserts, *International Journal of Refractory Metals & Hard Materials* 27 (2009) 754-763
- [10] J. Kundrak, A. G. Mamalis, K. Gyani, V. Bana: Surface Layer Microhardness Changes with High-Speed Turning of Hardened Steels, *International Journal of Advanced Manufacturing Technology* 53 (1-4) (2011) 105-112
- [11] I. Sztankovics, J. Kundrak: Determination of the Chip Width and the Undeformed Chip Thickness in Rotational Turning, *Key Engineering Materials* 581 (2014) 131-136
- [12] I. Sztankovics, J. Kundrak: Effect of the Inclination Angle on the Defining Parameters of Chip Removal in Rotational Turning, *Manufacturing Technology* 14 (1) (2014) 97-104

- [13] R. S. Pawade, H. A. Sonawane, S. S. Joshi: An Analytical Model to Predict Specific Shear Energy in High-Speed Turning of Inconel 718, *International Journal of Machine Tools & Manufacture* 49 (2009) 979-990
- [14] D.-J. Waldorf, R.-E. DeVor, S.-G. Kapoor: Slip-Line Field for Ploughing during Orthogonal Cutting, *Journal of Manufacturing Science and Engineering* 120 (4) (1998) 693-698
- [15] D.-J. Waldorf, R.-E. DeVor, S.-G. Kapoor: 1999, An Evaluation of Ploughing Models for Orthogonal Machining, *Journal of Manufacturing Science and Engineering* 121 (1999) 550-558
- [16] M. Annoni, G. Biella, L. Rebaioli, Q. Semeraro: Microcutting Force Prediction by Means of a Slip-Line Field Force Model, *Procedia CIRP* 8 (2013) 558-563
- [17] V. N. Gaitonde, S. R. Karnik, J. Paulo Davim: Selection of Optimal MQL and Cutting Conditions for Enhancing Machinability in Turning of Brass, *Journal of Materials Processing Technology* 204 (2008) 459-464
- [18] W. Zębala, R. Kowalczyk: Estimating the Effect of Cutting Data on Surface Roughness and Cutting Force during WC-Co Turning with PCD Tool using Taguchi Design and ANOVA Analysis *The International Journal of Advanced Manufacturing Technology* (2014) 1-16
- [19] I. Hanafi, A. Khamlichi, F. M. Cabrera, P. J. N. Lopez, A. Jabbouri: Fuzzy Rule based Predictive Model for Cutting Force in Turning of Reinforced PEEK Composite, *Measurement* 45 (2012) 1424-1435
- [20] C. Fetecau, F. Stan: Study of Cutting Force and Surface Roughness in the Turning of Polytetrafluoroethylene Composites with a Polycrystalline Diamond Tool, *Measurement* 45 (2012) 1367-1379
- [21] B. de Agustina, C. Bernal, A. M. Camacho, E. M. Rubio: Experimental Analysis of the Cutting Forces Obtained in Dry Turning Processes of UNS A97075 Aluminium Alloys, *Procedia Engineering* 63 (2013) 694-699
- [22] H. Joardar, N. S. Das, G. Sutradhar, S. Singh: Application of Response Surface Methodology for Determining Cutting Force Model in Turning of LM6/SiCP Metal Matrix Composite, *Measurement* 47 (2014) 452-464
- [23] R. Horvath, Á. Drégelyi-Kiss, Gy. Mátyási: Application of RSM Method for the Examination of Diamond Tools, *Acta Polytechnica Hungarica* 11:(2) (2014) 137-147
- [24] R. Horváth, Gy. Mátyási, Á. Drégelyi-Kiss: Optimization of Machining Parameters for Fine Turning Operations based on the Response Surface Method, *ANZIAM Journal* 55 (2014) C250-C265
- [25] R. Horváth, Á. Drégelyi-Kiss: Analysis of Surface Roughness of Aluminium Alloys Fine Turned: United Phenomenological Models and Multi-Performance Optimization, *Measurement* 65 (2015) 181-192

- [26] R. Horváth, Á. Drégelyi-Kiss: Analysis of Surface Roughness Parameters in Aluminium Fine Turning with Diamond Tool, Measurement 2013 Conference, Smolenice, Slovakia, 275-278
- [27] R. Horváth, Á. Czifra, Á. Drégelyi-Kiss: Effect of Conventional and Non-Conventional Tool Geometries to Skewness and Kurtosis of Surface Roughness in Case of Fine Turning of Aluminium Alloys with Diamond Tools, The International Journal of Advanced Manufacturing Technology (2014) 1-8
- [28] O. Kienzle, H. Victor: Die bestimmung von kräften und leistungen an spanenden werkzeugmaschinen, VDI-Z 94 (1952) 299-305
- [29] S. R. Frey: Repetitorium: Spanende Formung Schweizer Maschinenmarkt, 32 (1980) 26-29
- [30] C. Bus, N. A. L. Touwen, P. C. Veenstra, A. C. H. Van Der Wolf: On the Significance of Equivalent Chip Thickness, Annals of the CIRP, XXIV. (1971) 121-124
- [31] R. Horváth, S. Sipos, Gy. Mátyási: A New Model for Fine Turning Forces (in Hungarian) GÉP 6-7 (2014) 50-55

A CERIF Compatible CRIS-UNS Model Extension for Assessment of Conference Papers

**Siniša Nikolić*, Zora Konjović*, Valentin Penca*, Dragan
Ivanović*, Dušan Surla****

*University of Novi Sad, Faculty of Technical Sciences

Trg Dositeja Obradovića 6, 21000 Novi Sad, Serbia, sinisa_nikolic@uns.ac.rs,
ftn_zora@uns.ac.rs, valentin_penca@uns.ac.rs, chenejac@uns.ac.rs

**University of Novi Sad, Faculty of Sciences

Trg Dositeja Obradovića 3, 21000 Novi Sad, Serbia, surla@uns.ac.rs

Abstract: This paper proposes an extension to CERIF compatible CRIS, enabling automated evaluation of research achievements by applying diverse (country/region/institution specific), or even multiple evaluation rulebooks and guidelines. It was implemented as an extension to the CERIF compliant CRIS system of the University of Novi Sad (CRIS UNS) that already contains information for assessment of results from scientific journals, so the focus of this research is an extension to the CERIF model aimed at evaluation of the results published through conferences. Based on a survey and an analysis of selected evaluation rulebooks and guidelines, the paper proposes an extended CERIF model that comprises conference evaluation related metadata and a machine-readable representation of a rulebook that enables automated evaluation. A rule-based expert system is proposed for representation of evaluation rules and evaluation of research results. The Serbian rulebook is represented and implemented using the expert system Jess in order to evaluate the proposed model. Reliance of the model on CERIF standard allows its easy application in any CERIF compatible CRIS system.

Keywords: CERIF; model extension; conferences; automated evaluation; Jess

1 Introduction

In recent years, as research and innovations are becoming crucially important for economic development and Government financing is tightening, assessment of research achievements and capacities become an unavoidable condition for identifying high quality research, both for strategic planning and other purposes, like appointments to scientific/teaching positions, ranking researchers and/or research institutions, decisions on scientific project financing, etc. [1]. The Committee for Evaluation of Research defines research evaluation as a process based on critical analysis of data which leads to a judgment of merit [2]. Research

outputs, such as monographs/books, journal articles and conference papers are the evidence of a research study findings and they are the most suitable for assessment. Evaluation of conference proceedings papers is not as extensive as the evaluation of journal articles. This could be explained by the well-established opinion that journals present scientific results of the highest quality, but that cannot be an excuse, as conference papers have a capacity to exceed journal papers regarding the up-to-date presentation of ideas. This is due to a relatively short review time for scientific conferences. Also, for some fast-growing scientific areas (e.g. Computer Science - CS), conference papers are the major form of publishing (conference-cantered publication culture) [3]. As the papers published at conferences are a significant part of scientific production, it is necessary to investigate the assessment of those research results as well.

Electronic databases containing research outputs were, and still are, a basis for research evaluation process. Construction of an information system is necessary for efficient evaluation of scientific-research data [4]. The Current Research Information System (CRIS) that is based on the Common European Research Information Format (CERIF) standard is an example of such a system. It represents a good base for development of a system for evaluation of scientific results.

Usually, the evaluation process is carried out by a commission of domain experts that decide on huge amounts of publications by following some evaluation rulebook or set of guidelines that could be subject to change and/or subjective interpretation. Therefore, the evaluation should be supported by an evaluation engine that will apply rules automatically, provide explanations of the evaluation process, and even support an option for rewriting rules, if necessary. An expert system in which evaluation rules are expressed by some declarative language could be a solution to the problem.

2 Related Work

The fundamental concepts underlying the research presented in this paper represent an approach to evaluation (objects of evaluation, procedures and metrics applied to evaluate objects, entities that carry out evaluation) and information resources and software tools that are used to assist evaluation.

In practice, research performance evaluation, mainly relies on the quality of the resulting publications, i.e. scientific publications are primary objects of evaluation. Other objects of evaluation (e.g. researchers, research institutions, etc.) are evaluated mainly relying upon evaluations of the corresponding primary objects. As stated in [5], the quality of a publication can be determined by using different approaches: expert opinion (peer review), bibliometric evaluation, or a combination of these two (experts that use bibliometric data for their decisions,

the most acceptable approach to evaluation of research results so far). The authors of the paper [6] state that rankings have become one of the main forms of quality assessment in higher education over the past few decades. Publication rankings are based on individual (each publication gains individual score) or collective (all publications gain the same score as a part of a larger publication, with an optional coarse granulation like *scientific paper*, *professional paper*) evaluation of publications. Many evaluation principles are based on collective evaluation, arguing that assessment of the source (e.g. journal, conference) in which the article is published is unbiased, less time- and resource-consuming, and more economical than the individual publication evaluation [7]. So, the "quality/reputation" of a conference can be a dominant criterion when assessing the quality of conference results.

One of the most successful attempts to rank and evaluate conferences, so far, is a subjective CORE ranking [8]. Examples of conference rankings that are based on a voting procedure (classification) are ERA 2010¹ and Perfil-CC² Ranking for computer science area. The rankings mentioned above indicate that expert evaluation is only carried out for a relatively small number of conferences, usually for a particular scientific field and to satisfy the need of some country or geographical area. Evaluation on a larger scale requires use of some metric for conferences like acceptance rates, bibliometric indicators, and bibliometric related data [9].

Once decided on the assessment approach, it is necessary to have access to data required to carry out evaluation. Regarding this issue, it is important to mention that there are scientific publication databases, which store some metrics that can be used to evaluate conferences. Some of those databases are *Google Scholar Metrics*³ (GSM), *Microsoft Academic*⁴ (MA), *Web of Science Conference Proceedings Citation Index*⁵ (CPCI), *Elsevier Scopus*⁶ (ES) and *CiteSeer Estimated Venue Impact Factors*⁷ (EVIF). GSM, MA, ES and CPCI contain bibliography and citation data for proceedings articles. MA provides ranking of conferences, while GSM and EVIF provide single list ranking of journals and conferences, meaning that for ranking of conferences only journals must be somehow excluded. GSM, MA, and EVIF have free access; while ES and CPCI are commercially available products (access is for commercial subscribers only). CPCI contains the highest amount of data compared to other databases and can be indirectly used to create conference evaluation metrics. Considering domain

¹ <http://lamp.infosys.deakin.edu.au/era/?page=cforsel10>

² <http://www.latin.dcc.ufmg.br/perfilccranking/>

³ https://scholar.google.com/citations?view_op=top_venues&hl=en

⁴ <http://academic.research.microsoft.com/>

⁵ <http://thomsonreuters.com/conference-proceedings-citation-index/>

⁶ <http://www.scopus.com/>

⁷ <http://citeseerx.ist.psu.edu/stats/venues>

coverage and usability, these databases are severely limited in research scope and data completeness. They cover mostly the area of Computer Science and sometimes Electrical Engineering, and even for those, the data is not as complete as for journal articles.

In addition to the abovementioned data sources, there are databases and repositories of research institutions and/or states which are considerable and valuable information resources for conference results evaluation. Numerous institutional/state databases/repositories provide data in accordance with CERIF model. CERIF is an open [10], widespread [11] international standard, with an already proven capacity for research results evaluation [12, 13] and clearly recognized by European Science Foundation for that purpose [14]. An important advantage of CERIF is that it can be extended and adapted to different needs. This has already been proven through results, such as those aimed at storing bibliometric indicators [15], and an extension aimed at evaluation of journal papers in CRIS UNS, presented in [16]. The latest activities, aimed at CERIF integration, with complementing standards aimed to support effective evidence-based institutional decision-making are in progress [17]. All this makes CERIF an important reference point for the purpose of evaluation regulated by national rulebooks.

In order to enable efficient evaluation, it is necessary to provide evaluators with tools that will automate evaluation process as much as possible. In an evaluation of research results, where rules from rulebooks are applied by some commissions, an obvious approach to evaluation automation is to use some rule-based expert system. Unlike the conventional systems that solve a problem by executing a well-defined algorithm, expert systems rely upon a knowledge base that contains statements and facts about the problem [18]. The advantage of a rule-based (declarative programming) code over a conventional programming code is that it is much easier to read, maintain and change, even by those who are not familiar with programming [19].

3 Analysis of National Rulebooks and Guidelines

The purpose of this section is to analyze evaluations at the national levels that apply to papers that are presented at conferences. In general, evaluations at national levels differ from one country to another. For this analysis, we have chosen representative countries differing in size, economic development, geo-location, and relation to the EU. In the rest of this section, country specific rulebooks and/or guidelines are analyzed for four non EU - Southern European countries (Serbia, Macedonia, Montenegro, Bosnia and Herzegovina), five EU countries (Croatia, Slovenia, Czech Republic, Hungary, United Kingdom), and one non-European country (Australia).

The evaluation of research results in Serbia utilizes the "combined" approach. For evaluation purposes, commissions use the document that was prescribed by the Serbian Ministry of Education and Science⁸. The results presented at conferences are evaluated based on the category of a conference and type of presentation and result. Conferences are categorized as international or national by the commission that is in charge of a scientific field corresponding to the scope of a conference.

The rulebook of Bosnia and Herzegovina is identical to the Serbian rulebook in terms of data requirements⁹. It prescribes the same categorization as the Serbian rulebook, differing only in categorization codes.

In Macedonia, the assessment of researchers' results is done by a commission. The commission uses the document that is created by the leading Macedonian university - University Ss. Cyril and Methodius in Skopje¹⁰. Conference results are categorized on the basis of conference evaluation. There are two types of conferences: a scientific/expert conference with an international program committee and a scientific/expert conference. Conferences that have an international committee with members from at least 5 countries, where the number of members from the most represented country does not exceed 40%, are ranked as higher quality conferences.

In Montenegro, the evaluation approach is almost the same as in Serbia and Macedonia (done by a commission and relying on bibliometric data). A sort of specific quality of the Montenegrin rulebook¹¹ is that the information on editors is used for conference proceedings results evaluation. Conference results are categorized on the basis of evaluation of a conference. The evaluation of a conference relies on: proceedings publication language (worldwide accepted languages are favored), proceedings editorial committee structure (international committees with distinguished members are favored) and organizer (international organizers are favored).

The Croatian evaluation rulebook¹² and its rules for classifying conference results are based on the evaluation of a conference. The conference category is determined based on the organizer (organizer must be a part of a specially verified list - organizer type) and the indexing of conference proceeding in CPCI (proceeding must be a part of list). A conference is categorized as international if organized by an international scientific association or if its proceeding is indexed in CPCI.

⁸ http://www.mpn.gov.rs/images/content/nauka/pravna_akta/PRAVILNIK_o_zvanjima.pdf

⁹ http://mcp.gov.ba/org_jedinice/sektor_nauka_kultura/pravni_okvir/podzakonski_akti/default.aspx?id=3379&langTag=bs-BA

¹⁰ http://www.ffk.pesh.mk/Vazni_dokumenti/Pravni_dokumenti/22_133786659.pdf

¹¹ http://www.ucg.ac.me/zakti/akademsko_zvanja.pdf

¹² http://narodne-novine.nn.hr/clanci/sluzbeni/2013_03_26_447.html

The Slovenian rulebook¹³ assesses researchers' results based on a set of criteria for research excellence (evaluation of researchers' publications and citations). Scientific contributions from conferences are categorized by the evaluation of a conference. By the Slovenian rulebook, conferences that take place abroad or those organized in Slovenia in worldwide accepted languages, with an international committee and a minimum of one-third of publications from abroad gain the highest score. The particulars of the Slovenian rulebook are two predefined time frames for citation (last 5 and 10 years). The total number of citations and the number of pure citations (citations without auto-citations) in the CPCI for the defined time frames are required for conference papers.

The Czech rulebook¹⁴ evaluates each submitted publication by using bibliometric data or scientific area experts within panels. Publications from conferences are not preferable in some scientific areas (not accepted, or included in a share of 50%, 25%, or below 5% of all publications). Conference papers are evaluated directly by experts in panels (that approach does not require evaluation of the conference itself). A conference proceedings paper is eligible for evaluation only if over 2 pages in length and if the conference proceeding in which the paper is published is indexed in CPCI or Scopus (proceeding must be a part of CPCI or Scopus list).

PhD and Habilitation committees at the Faculty of Science and Information Technology in Szeged, Hungary, directly evaluate conference papers. The Document¹⁵ is used as a guideline for that evaluation. What is particular for this evaluation compared to the others is that it accounts for pure citations which are not derived from the author's affiliation, pure citations from any source and pure citations that are only found in a predefined list of databases (e.g. Web of Science and Scopus). For a conference paper to be considered for evaluation, its conference proceeding must be indexed in a predefined list of databases (e.g. Mathematical Reviews, WoS, CPCI, Scopus etc.) or its conference must be an internationally recognized event (symposium and workshop) with an acceptance rate below 50% and the ratio of international authors above 50%. According to commissions, a conference is international if the participants and the committee are international. Since those conditions are not formally supported with some metrics, such as the numbers of authors and committee members (no such details are used, as *Committee members structure* and *Results data* are used in other rulebooks), there are no formal criteria for representing them. The quality of each conference paper is determined based on the opinion of commissions, categorizing them on a 3 grade scale as International in English, Domestic in English or In Hungarian.

¹³ <http://www.arrs.gov.si/en/akti/prav-sof-ocen-sprem-razisk-dej-sept-11.asp>

¹⁴ <http://www.vyzkum.cz/FrontClanek.aspx?idsekce=695512&ad=1&attid=695694>

¹⁵ <http://www.sci.u-szeged.hu/kar/kari-szabalyzatok/ttik-doktori-szabalyzat?objectParentFolderId=19613>

The United Kingdom has created an assessment framework called REF¹⁶ (Research Excellence Framework). According to REF, conference papers are directly evaluated by experts on Units of Assessments (panels). A particular quality of REF is that it requires data about the abstract of publication. In some cases, citation data might be utilized in evaluation (for a particular scientific area, if it is available and if experts would like to use it). REF states that all research output data requirements are compatible with CERIF.

ERA¹⁷ (Excellence in Research for Australia) assessment framework in Australia is very close to REF, i.e. conference papers are evaluated by experts in panels. Some panels accept peer review evaluation, but only for up to 30% of the submitted publications, so conference proceedings papers are only evaluated by this approach.

3.1 Metadata for Evaluation of Conference Papers

The analysis of rulebooks and guidelines leads to the conclusion that all national rulebooks somehow include conference papers. In general, scores are assigned to conference proceeding articles either based on the opinion of the experts in the panel (experts' judgment which can, up to a certain extent, rely upon bibliometric indicators), or by a commission applying the rules, relying upon conference categories, where a conference category is determined based on common criteria (language, structure of a scientific committee, etc.), and, sometimes, on bibliometric indicators that apply to conference proceedings (CPCI indexing). In some rulebooks, when assessing publications presented at conferences, in addition to the metadata describing the publication itself, it is necessary to include the data for the conference and conference proceedings. Therefore, metadata consists of three sets: conference metadata, proceedings metadata and publication metadata.

As a result of the rulebooks' analysis, the following is a set of evaluation metadata for conferences (Table 1) that all rulebooks include: *conference name*, *year*, *place*, *presentation language*, *proceeding data*, *conference committee structure*, *organizer data* and *conference results data*. In most of the rulebooks, the category of a conference depends on the conference organizer, conference language, proceeding data and conference committee structure. The conference committee structure metadata includes the *total number of committee members*, *total number of countries from which committee members originate*, *number of committee members from the most represented country*. The metadata on the conference organizer describe the particular organizer (*name of institution*) and *organizer type* (international, national). The conference results data is described with the *total*

¹⁶ http://www.ref.ac.uk/media/ref/content/pub/panelcriteriaandworkingmethods/01_12.pdf

¹⁷ http://www.arc.gov.au/pdf/era12/ERA%202012%20Evaluation%20Handbook_final%20or%20web_protected.pdf

number of papers, total number of submitted papers and number of papers whose authors are foreigners.

Table 1
Evaluation metadata for conferences

	Serbia	Macedonia	Montenegro	Bosnia and Herzegovina	Croatia	Slovenia	Czech Republic	Hungary	United Kingdom	Australia
name	+	+	+	+	+	+	+	+	+	+
year	+	+	+	+	+	+	+			+
place	+	+	+	+	+	+	+			+
organizer	+		+	+	+					
presentation language	+			+		+				
proceeding	+	+	+	+						
conference committee structure: total number of committee members	+	+		+		+				
conference committee structure: total number of countries from which the committee members originate	+	+		+		+				
conference committee structure: number of committee members from the most represented country		+								
organizer data: organizer type	+		+	+	+					
conference results data: total number of papers	+			+		+		+		
conference results data: number of submitted papers								+		
conference results data: number of papers whose authors are foreigners	+			+		+		+		

Conference proceeding (Table 2) may include data as proceedings: *title, ISBN, publication language, publisher, editorial committee structure* and *indexing of the conference proceeding in CPCI and/or SCOPUS*. Publication language, indexing and editorial committee structure are used to categorize a conference. The editorial committee data, which is used only by the Montenegro rulebook, consists of *name(s) of editor(s), total number of editorial committee members* and *total number of countries from which the committee members originate*. Indexing of conference proceedings in CPCI and/or SCOPUS is required only for the Croatian and Czech Republic rulebooks. In Hungary, it is preferred that the conference proceeding be indexed in a predefined list of databases. The metadata describing the publisher are *publisher's name* and *headquarters of the publisher*.

A conference publication (Table 3) is defined by the following data: *title, author(s) name(s), publication year, total number of pages, conference data, proceedings data, citation data, DOI, URL, scientific area, abstract* and *type of evaluation entity*. The citation data is required only in the UK (*total number of citations*), Slovenia (*total number of citations, number of pure citations in the last 5 and number of pure citations in the last 10 years*) and Hungary (*total number of citations, number of pure citations, number of pure citations not derived from the author's own affiliation, number of pure citations found in a predefined list of databases*).

Table 2
Evaluation metadata for conference proceedings

	Serbia	Macedonia	Montenegro	Bosnia and Herzegovina	Croatia	Slovenia	Czech Republic	Hungary	United Kingdom	Australia
title	+	+	+	+	+	+	+	+	+	+
ISBN			+				+	+	+	
editor(s)			+							
publisher			+							+
publishers' data: headquarter of the publisher										+
publication language	+		+	+						
indexing of conference proceeding: in CPCI					+		+	+		
indexing of conference proceeding: in SCOPUS							+	+		
indexing of conference proceeding: in a predefined list of databases								+		
editorial committee structure: total number editorial committee members			+							
editorial committee structure: total number of countries from which the committee members originate			+							

Table 3
Evaluation metadata for conference publication

	Serbia	Macedonia	Montenegro	Bosnia and Herzegovina	Croatia	Slovenia	Czech Republic	Hungary	United Kingdom	Australia
title	+	+	+	+	+	+	+	+	+	+
author(s)	+	+	+	+	+	+	+	+	+	+
publication year	+	+	+	+	+	+	+	+	+	+
total number of pages	+	+	+	+	+	+	+	+	+	+
conference	+	+	+	+	+			+		+
conference proceeding			+		+			+		+
DOI/URL							+		+	+
scientific area						+	+	+	+	+
abstract									+	
type of evaluation entity	+	+	+	+	+	+				
citation data: total number of citations						+		+	+	
citation data: number of pure citation								+		
citation data: number of pure citation in the last 5						+				
citation data: number of pure citation in the last 10						+				
citation data: number of pure citation not derived from the authors own affiliation								+		
citation data: number of pure citation found in predefined list of databases								+		

4 Conference Evaluation-related Data in the CERIF Model

CRIS is an information system that stores bibliographic and normative data for entities related to conference results, as well as data on their interrelations. Because the CERIF standard and its physical model are a basis of CRIS systems, a legitimate conclusion is that the CERIF model should be investigated against the existing support of the relevant evaluation data, as well as against the possible extensions if needed.

With CERIF, it is possible to determine the following entities: people involved in research activities (e.g. authors, organization members, etc.), organizations (e.g. universities, government agencies, publishing houses, etc.), research projects, research results (scientific publications, patents, etc.), etc. There are 8 main groups of CERIF entities. *Base entities* represent the core (basic) model entities (cfPers, cfOrgUnit and cfProject). *Result entities* represent the results from scientific research (cfResPubl, cfResProd and cfResPat). *Infrastructure entities* represent the infrastructure that is relevant for scientific research (like cfEquip, cfFacil etc.). *2nd Level entities* further describe the *Base* and *Result Entities* (e.g. cfEvent, etc.). *Indicator and Measure entities* are used to define the research impact and supporting claims of that impact, covering the abovementioned entities (cfIndic and cfMeas). *Multiple language entities* provide multilingualism for CERIF data, like cfResPublTitle. *Semantic layer entities* cfClass (classes) and cfClassScheme (classification schemes) enable a rich semantic representation of data. CERIF prescribes a vocabulary that might be utilized for establishing classification, e.g. class "Author" of scheme "Person Output Contributions" can be used to define the person that is the author of a conference paper. *Link entities* are used to state time-determined relations among other entities, like relation of a person and a publication cfPers_ResPubl. Every Link entity is described with a role (cfClassId, cfClassSchemeId), timeframe of relation (cfStartDate, cfEndDate), value (cfFraction) and identifiers of elements creating the relation (e.g. cfPersId, cfResPublId). The "role" in link entities is not stored directly as an attribute value, but as references to Semantic layer entities.

4.1 The Data in the CERIF Model

The CERIF model Version 1.5 is used in this paper as a basis for proposing a model for evaluation of the results from conferences. The existing CERIF model for storing proceedings, conference results and conference data entities is represented in Figure 1. For simplicity and readability of the diagram, the classification entity *cfClass* was omitted.

The data about proceedings and conference results in CERIF can be placed in entities *cfResPubl*, *cfResPublTitle* and *cfResPublAbstr*. The title of a paper is acquired from *cfResPublTitle*, while the abstract of a proceeding paper is acquired from the entity *cfResPublAbstr*. The type of publication is set up by classification of instances of *cfResPubl* via *cfResPubl_Class*. CERIF scheme *Output Types* and its classifications from the controlled vocabulary (e.g. *Conference Proceedings*, *Conference Proceedings Article*, *Conference Poster*, *Conference Abstract*, *Conference Contribution*, etc.) are used for that.

Conference data can be stored in the event entity *cfEvent* and its name *cfEventName*. The event is stated as a conference via the *cfEvent_Class* entity and by the CERIF scheme *Event Types* and class *Conference*. The links between instances of publications and events are saved in the entity *cfResPubl_Event*.

The researchers that are authors, editors or reviewers of a publication can be represented by instances of entities *cfPers*, *cfPersName* and *cfPersName_Pers*, which store information about a person and persons' name. An instance of *cfPers* is connected to instance *cfResPubl* with the link entity *cfPers_ResPubl*.

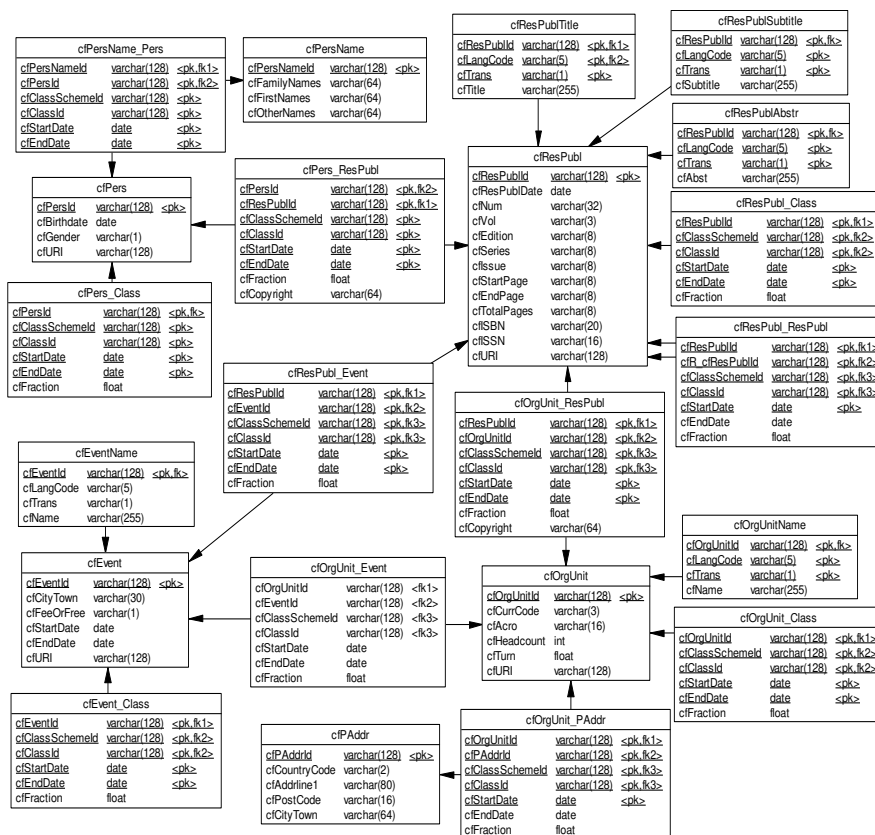


Figure 1

Existing CERIF entities for storing proceedings, conference results and conference data

The publisher of proceedings or the organizer of a conference can be represented by instances of entities *cfOrgUnit* (information about organizations) and *cfOrgUnitName* (organizations' names). The information about the headquarters (place) of an organization is acquired from the attributes *cfCountryCode* and *cfCityTown* of the entity *cfPAddr*, which is connected to the organization through an instance of *cfOrgUnit_PAddr* entity. Organizations are linked to publications and events with link entities *cfOrgUnit_ResPubl* and *cfOrgUnit_Event* respectively. A role of an organization that is a publisher (of proceeding) is enabled with CERIF scheme *Organisation Output Roles* and class *Publisher*. The role of a conference organizer is defined by using CERIF scheme *Organisation Output Contributions* and class *Host*.

4.2 CERIF Extension for Evaluation of Conference Publications

CRIS UNS is a CERIF compatible research management system that has been developing since 2008 at the University of Novi Sad in the Republic of Serbia¹⁸. An extension to CERIF model which incorporates metadata for evaluation of journal articles in CRIS UNS was proposed in [20], while an extension to CERIF model aimed at modelling and storing a rulebook was proposed in [21]. The previously proposed CERIF model extension related to rulebook representation enables representation of rulebooks that relies on classifications.

An extension to CERIF aimed at providing the data necessary for evaluation in accordance with the findings from Section 3 is accomplished by relying on (using "as is" or repurposing) some of the existing CERIF entity attributes and by adding new semantics (classifications schemes and classes) for the existing CERIF entities. The semantics is defined to comply with the analyzed rulebooks. For the purpose of enabling a relation between a complex result and its constituents, a new class *Belongs to* and the corresponding scheme *General Relations* are added to CERIF vocabulary. The class *Belongs to* is used to classify entities *cfResPubl_ResPubl* (stating the inner relations among publications) and *cfResPubl_Event* (stating the relation between proceeding/paper and conference).

cfResPubl attributes (Figure 1) provide evaluation information for publication year (*cfResPublDate*), ISBN or ISMN identifier (*cfISBN*), DOI or URL identifier (*cfURI*) and the total number of publication pages (*cfTotalPages*). The publication language is acquired from the entity *cfResPublTitle*. Assuming that the original language of the title is the same as the language of the publication, the attribute *cfLangCode* can be interpreted as publication language. The scientific area (group) for proceedings papers is provided via an instance of entity *cfResPubl_Class*, by stating the relation with the appropriate scientific area/group (the classification of scientific area/group is already defined within the rulebook). The indexing of a conference proceeding in CPCI or SCOPUS is enabled through instances of the entity *cfResPubl_Class*. The classification is enabled by a new scheme *Proceeding Evaluation Details* and three new classifications *Is indexed in CPCI*, *Is indexed in SCOPUS* and *Is indexed in a predefined list of databases*.

The year and place of a conference are obtained from *cfEvent* attributes containing data on timeframe (*cfStartDate*, *cfEndDate*), as well as the country and city (*cfCountryCode*, *cfCityTown*) of the conference. Assuming the same for *cfEventName* as for *cfResPublTitle*, the attribute *cfLangCode* of *cfEventName* can be utilized to acquire conference presentation language.

¹⁸ <http://www.cris.uns.ac.rs>

The organizer type for a conference is enabled through the entity *cfOrgUnit_Class* that is used for classifying organization (*cfOrgUnit*) as international/national organizer. The categorization of organizer is done by new scheme *Organiser Evaluation Details* and its classes *International Organiser* and *National Organiser*.

The data necessary for evaluation, which is measurable (e.g. total number of committee members, total number of papers, total number of citations etc.), can be stored by relying on CERIF Indicator and Measure entities (Figure 2). The authors of [17] used those entities to store metrics for comparison of universities (research inputs, process, outputs and outcomes). With *cfIndic* and *cfMeas* it is possible to create various quality and quantity indicators which can be related to the base, result and 2nd level entities.

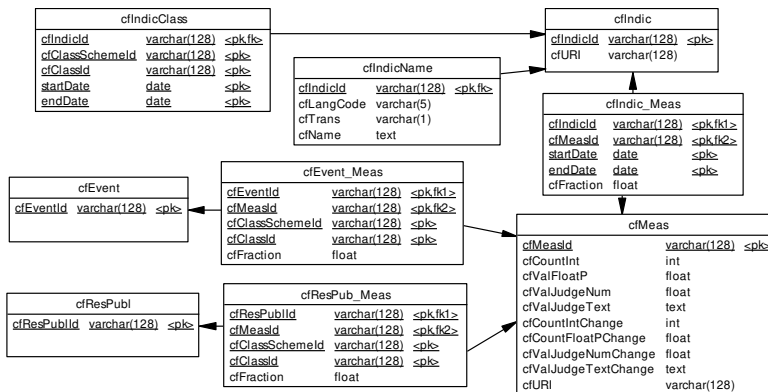


Figure 2

Storing measurements for publications and events

Every measurement identified in section 3 is represented as an instance of *cfIndic* (semantics of the measured value). The name of the measurement is kept in the multilingual entity *cfIndicName*. The classification of instances is done via the entity *cfIndic_Class* and a newly created scheme *Evaluation Indicator Measurement* and its classes. The class *Citation data* classifies the *cfIndic* instances Total number of citations, Number of pure citations, Number of pure citations in the last 5 years, Number of pure citations in the last 10 years, Number of pure citations not derived from the author's own affiliation and Number of pure citations found in a predefined list of databases. The classification *Conference committee structure* is used for *cfIndic* instances named Total number of committee members, Total number of countries from which committee members originate and Number of committee members from the most represented country. The class *Conference results data* is applied to *cfIndic* instances named Total number of papers, Number of submitted papers and Number of papers whose authors are foreigners. The class *Editorial committee structure* classifies the *cfIndic* instances Total number of committee members and Total number of

countries from which the committee members originate. The concrete value of measurements for a single publication or a single event are defined as instances of *cfMeas*, storing integer (*cfCountInt*), float (*cfValFloatP*) or textual data (*cfJudgeText*) corresponding to the measured value. Measurements are connected by links (*cfResPubl_Meas*, *cfEvent_Meas*) to entities that are characterized by measurements (e.g. *cfEvent_Meas* connects an instance *cfMeas* storing Total number of committee members with a *cfEvent* instance representing a conference). The classification of those links is done with the class *Belongs To* from the scheme *General Relations*. The same class is used to link measurement and indicator (e.g. *cfMeas* is linked to *cfIndic* by the link entity *cfIndic_Meas*).

5 Validation: The Case Study of the Serbian Rulebook

In the Serbian rulebook, the Ministry prescribes the classification (categorization) of research results for Serbian researchers. The Serbian rulebook prescribes classification (categorization) of research results organizing them into hierarchical levels. Conferences can be categorized as international (M30) or national (M60) scientific meetings. So, a scientific result belonging to a M30 conference can be categorized as *invited talk paper printed in full* (M31), *invited talk paper printed as abstract* (M31), *paper printed in full* (M33), *paper printed as abstract* (M34), *authorized discussion paper* (M35), and *editorial of proceedings* (M36). The categorization of scientific results belonging to an M60 conference is similar to the one used for an M30 conference.

In order to provide a flexible and efficient mechanism for representation of machine-readable rulebooks aimed at automated evaluation of conference papers, we propose a solution relying on a rule-based expert system. For the purpose of representing rulebooks and evaluation rules, a Jess (Java Expert System Shell¹⁹) rule-based system is selected. Jess programming language is a Lisp-like declarative rule-based language that is very easy to read and understand by non-programmers, which was the main reason for our commitment to Jess as a solution to a rulebook representation task.

5.1 Jess Implementation of the Serbian Rulebook

The Serbian rulebook and its classification scheme are represented in Jess as facts. All fact templates are derived from Java object representation (Java Bean Classes) of the proposed CERIF extension. An example of a Jess code illustrating that

¹⁹ <http://herzberg.ca.sandia.gov/>

principle can be "*deftemplate CfMeas (declare (from-class CfMeas))*", where a Jess shadow template *CfMeas* is created by looking at Java class representing CERIF *CfMeas* entity. The primitive Java class properties (e.g. *cfCountInt* of type integer) are mapped by default to Jess template slots with the same name and the same/similar type (e.g. Jess INTEGER). Non-primitive Java properties (if any) are mapped to Jess *OBJECT* type slots with the same name as corresponding Java properties. Since the slots of type *OBJECT* in Jess hold a reference to the Java object itself, the original objects from Java are always easily accessible.

The data for the facts representing the Serbian rulebook is extracted from Java object instances of extended CERIF model entities *RuleBook*, *RuleBook_Class*, *RuleBookName*, *RuleBookDescr*, *RuleBook_ResearchersRole*, *RuleBook_ResultsType*, *RuleBook_EntityType*, *ResultsTypeMeasure*, whose modelling is presented in [21].

Every evaluation rule from the Serbian rulebook can be formulated as a Jess rule with a distinct priority, where LHS (Left Hand Side) is composed of the fact statements that are all connected with logical conjunction by default. This priority enables avoidance of multiple classifications (e.g. higher priority is assigned to the rules classifying a result as "international" rather than the ones classifying it as "national"). All rules (Jess .clp files) for evaluation of conference results are available at http://s000.tinyupload.com/?file_id=95392329127505429716.

The proposed concept *CERIF model extension for storing evaluation data for conference papers* and the use of Jess rules language and reasoning engine for assigning categories to both conference and conference result (conference paper) are evaluated and verified by assessing the conference paper:

Nikolic, S., Penca, V., Ivanovic, D. (2014) "System for Modelling Rulebooks for the Evaluation of Scientific-Research Results. Case Study: Serbian Rulebook", *Proceedings of the 4th International Conference on Information Society and Technology (ICIST 2014)* Society for Information Systems and Computer Networks, Kopaonik, Serbia, March 9-13, 2014, pp. 102-107

Following the evaluation process by the Serbian rulebook, the conference ICIST 2014 should first receive a category, before it is possible to categorize its papers. So, the higher priority rules for categorizing the conference are executed prior to the rules for categorizing papers. Also, in accordance with the Serbian rulebook, Jess will first try to classify the conference as international (M30) and, if that fails, it will attempt to apply other rules following the priority order (e.g. national conferences - M60, excluded from evaluation, etc.).

srRuleBook_international_conference (Figure 3) is a rule that classifies the conference as a M30 type.

The first paragraph in the rule checks if *4th International Conference on Information Society and Technology (ICIST 2014)* is a CERIF event (*CfEvent* with

variable *?cfEventId*) that is classified as a conference. *ICIST 2014* (*?cfEventId*) has a class (*cfOrgUnit_Class*) with attribute *cfClassId* having the value "Conference", so the condition "*{cfClassId == Conference}*" is met.

```
(defrule srRuleBook_international_conference ;rule head - name
;rule description
"Serbian rulebook categorisation of conference as an event of
international importance by organiser or conference committee and results data"
;rule priority
(declare (salience 900))

;first paragraph - checks satisfaction of basic requirements
?event <- (cfEvent (cfEventId ?cfEventId) {cfStartDate != nil} {cfEndDate != nil} {cfCountryCode != nil} {cfCityTown != nil})
(cfEvent_Class (cfEventId ?cfEventId) {cfClassSchemeId == "Event Types"} {cfClassId == "Conference"})
(cfEventName (cfEventId ?cfEventId) {cfName != nil } {cfLangCode ?presentationLang & (neq ?presentationLang nil)))

;second paragraph - match the rule if conference is not already evaluated
(not (cfEvent_CommissResultType (cfEventId ?cfEventId) {ruleBookId == "Serbian RuleBook"}))

;third paragraph - check existence of link between conference with its proceeding
(cfResPubl (cfResPublId ?cfResPublId))
(cfResPubl_Class (cfResPublId ?cfResPublId) {cfClassSchemeId == "Output Types"} {cfClassId == "Conference Proceedings"})
(cfResPublTitle (cfResPublId ?cfResPublId) (cfLang ?publicationLang & (neq ?publicationLang nil)))
(cfResPubl_Event (cfResPublId ?cfResPublId) (cfEventId ?cfEventId) {cfClassSchemeId == "General Relations"} {cfClassId == "Belongs To"})

;fourth paragraph - checks language conditions
(test (or (eq ?presentationLang "en") (eq ?presentationLang "fr") (eq ?presentationLang "de") (eq ?presentationLang "es")))
(test (or (eq ?publicationLang "en") (eq ?publicationLang "fr") (eq ?publicationLang "de") (eq ?publicationLang "es")))

;fifth paragraph - checks multiple conditions
(or
;first part - check organization conditions
(and (cfOrgUnit (cfOrgUnitId ?cfOrgUnitId))
(cfOrgUnit_Event (cfOrgUnitId ?cfOrgUnitId) (cfEventId ?cfEventId)
{cfClassSchemeId == "Organisation Output Contributions"} {cfClassId == "Host"})
(cfOrgUnit_Class (cfOrgUnitId ?cfOrgUnitId) {cfClassSchemeId == "CERIF Entities"} {cfClassId == "Organisation"})
(cfOrgUnit_Class (cfOrgUnitId ?cfOrgUnitId) {cfClassSchemeId == "Organiser Evaluation Details"} {cfClassId == "International Organiser"}))
;second part - checks commission structure and foreign authors papers conditions
(and
;total number of countries from which committee members originate condition
(cfIndic (cfIndicId ?cfIndicId1))
(cfIndic_Class (cfIndicId ?cfIndicId1) {cfClassSchemeId == "Evaluation Indicator Measurement"} {cfClassId == "Conference committee structure"})
(cfIndicName (cfIndicId ?cfIndicId1) {name == "Total number of countries from which committee members originate"})
(cfMeas (cfMeasId ?cfMeasId1) {cfCountInt >= 5})
(cfIndic_Meas (cfIndicId ?cfIndicId1) (cfMeasId ?cfMeasId1) {cfClassSchemeId == "General Relations"} {cfClassId == "Belongs To"})
(cfEvent_Meas (cfEventId ?cfEventId1) (cfMeasId ?cfMeasId1) {cfClassSchemeId == "General Relations"} {cfClassId == "Belongs To"})
;number of papers whose authors are foreigners condition
(cfIndic (cfIndicId ?cfIndicId2))
(cfIndic_Class (cfIndicId ?cfIndicId2) {cfClassSchemeId == "Evaluation Indicator Measurement"} {cfClassId == "Conference results data"})
(cfIndicName (cfIndicId ?cfIndicId2) {name == "Number of papers whose authors are foreigners"})
(cfMeas (cfMeasId ?cfMeasId2) {cfCountInt >= 10})
(cfIndic_Meas (cfIndicId ?cfIndicId2) (cfMeasId ?cfMeasId2) {cfClassSchemeId == "General Relations"} {cfClassId == "Belongs To"})
(cfEvent_Meas (cfEventId ?cfEventId2) (cfMeasId ?cfMeasId2) {cfClassSchemeId == "General Relations"} {cfClassId == "Belongs To"}))

;sixth paragraph - check total number of papers conditions
(cfIndic (cfIndicId ?cfIndicId3))
(cfIndic_Class (cfIndicId ?cfIndicId3) {cfClassSchemeId == "Evaluation Indicator Measurement"} {cfClassId == "Conference results data"})
(cfIndicName (cfIndicId ?cfIndicId3) {name == "Total number of papers"})
(cfMeas (cfMeasId ?cfMeasId3) {cfCountInt >= 10})
(cfIndic_Meas (cfIndicId ?cfIndicId3) (cfMeasId ?cfMeasId3) {cfClassSchemeId == "General Relations"} {cfClassId == "Belongs To"})
(cfEvent_Meas (cfEventId ?cfEventId3) (cfMeasId ?cfMeasId3) {cfClassSchemeId == "General Relations"} {cfClassId == "Belongs To"})
=>
(add (new cfEvent_CommissResultType ?event.cfEventId "Serbian RuleBook" "Results Type Scheme" "M90"
"Evaluated as International Conference by Serbian RuleBook")))
```

Figure 3

Rule for classifying CERIF events as M30 - conference of international importance

By the Serbian rulebook, an event must satisfy minimum requirements (e.g. data for *name*, *place*, *year*, etc. must be provided) to be considered for evaluation (first paragraph of the rule). *ICIST 2014* was organized from 09/03/2013 (value stored in attribute *cfStartDate*) to 13/03/2013 (value stored in attribute *cfEndDate*) at Kopaonik (value stored in attribute *cfCityTown*), Serbia (value stored in attribute *cfCountryCode*), so the conditions "*{cfStartDate != nil} {cfEndDate != nil} {cfCountryCode != nil} {cfCityTown != nil}*" are satisfied. The conference has a name (*cfEventName*) with attributes *cfName* and *cfLangCode*, where *cfEventName*

stores the value "*International Conference on Information Society and Technology (ICIST 2014)*" and *cfLangCode* stores the value "en" (conference presentation language). The variable *?presentationLang* holds the value that is acquired from *cfLangCode*. The attribute *cfName* and variable *?presentationLang* fulfil the conditions "*{cfName != nil}*" and "*(neq ?presentationLang nil)*". This means that all required data for a conference categorization is provided.

For an event to be an international conference, its presentation and proceeding publication languages must be worldwide accepted. The third paragraph checks if there is a connection between the conference and its proceedings. ICIST 2014 (*?cfEventId*) has proceedings *Proceedings of the 4th International Conference on Information Society and Technology* (*?cfResPubId*) whose title (*cfResPubName*) attribute *cfLangCode* stores value "en" (conference publication language). The variable *?publicationLang* holds the value that is acquired from *cfLangCode*. The condition "*(neq ?presentationLang nil)*" is met. The fourth paragraph checks the abovementioned language conditions. At ICIST 2014 the papers were presented and published in the English language, so language conditions "*(eq ?presentationLang "en")*" and "*(eq ?publicationLang "en")*" are fulfilled.

According to Serbian rulebook, international conferences are those whose organizer is an international scientific association/institution or whose committee and results have international characteristics (committee members are from at least 5 countries, the conference must have at least 10 papers whose authors are foreigners). That claim is defined by the fifth paragraph, where two logical statements (conjunctions) are connected via *OR* operator. The first conjunction statement checks the conditions for organization (*CfOrgUnit*). For ICIST 2014 (*?cfEventId*) there is an organizer *Information Society of Serbia* (*?cfOrgUnitId*) that is classified as "National Organiser" (value of attribute *cfClassId* of class *cfOrgUnit_Class*). So, the organizer condition "*(cfClassId == "International Organiser")*" is not met. The second conjunction checks if the conditions (values of measures) related to international characteristics of conference committee structure and conference results are met. The measures are defined as two individual indicators (*CfIndic*), each having a concrete measurement (*CfMeas*) that is linked to the event. Since ICIST 2014 conference committee members were from 17 countries, *Total number of countries from which the committee members originate* (*?cfIndicId1*) has a measurement *?cfMeasId1* with an attribute *cfCountInt* having the value 17, so the condition "*{cfCountInt >= 5}*" is satisfied. There were 29 papers whose authors are from foreign countries, so *Number of papers whose authors are foreigners* (*?cfIndicId2*) has a measurement *?cfMeasId2* with an attribute *cfCountInt* that has a value of 29, which means that the condition "*{cfCountInt >= 10}*" is met.

By the rulebook, the conference is excluded from evaluation if the total number of accepted papers is less than 10 (the sixth paragraph). Following the same principle as for the other explained measures, 85 papers from ICIST 2014 satisfy the last condition *Total number of papers*.

Having fulfilled the conditions *publication language*, *presentation language*, *conference committee structure* and *conference results data*, the conference is categorized as *M30*.

Once a conference receives a category, the conference paper "System for modelling rulebooks for the evaluation of scientific-research results. Case study: Serbian Rulebook" is assessed by the conference paper rules. The rule *srRuleBook_paper_printed_inFull_or_asAbstract* for categorizing M30 conference papers (*Conference Proceedings Article* or *Conference Poster*) is applied. That rule relies on the value *total number of pages* (attribute *cfTotalpages*) of CERIF publication entity *CfResPubl*, to determine the paper category (*paper published in full* - M33 or *paper published as abstract* - M34). That value is 6, so the paper is categorized as *M33*.

After the example mentioned above, we can conclude that by utilizing a rule-based expert system, it is feasible to evaluate research results in Serbia. By analogy, all other research results from the Serbian rulebook can be evaluated.

Conclusion

This paper investigates the current developments for evaluation of conference results in a broader research information environment that include various research evaluation data, with the primary aim to propose a data model and tools for automated evaluation of papers presented at scientific conferences.

In the first step, we have carried out and presented an analysis of selected national rulebooks regarding conference results evaluation. Then we proposed an extension to the CERIF model that supports evaluations in line with the results of the analysis. For evaluation of conference results, we have confirmed that all evaluation data required by the analyzed rulebooks are provided by the proposed CERIF model extension. By utilizing that extension, which can be applied to any CERIF like CRIS system, the evaluations on institution or national level should be easily introduced to research information systems.

For the automation of the evaluation process we propose an approach that utilizes the Jess rule language for representing a rulebook and its evaluation rules, and Jess inference engine for automated evaluation. The proposed approach was verified through an example of the Serbian rulebook and evaluation of a paper published at an international conference. Following that same principle, it can be shown that evaluation for other rulebooks and commissions can be accomplished by the same approach. This result could be useful for evaluation commissions, hence only a facts and rules should be specified, while all the evaluation is done by the inference engine.

The application of the proposed approach has certain constraints, regarding both the proposed CERIF extension and automation, of the evaluation process. Regarding the CERIF extension/model, the constraint is its support, since it analyses only current versions of the selected rulebooks. This, together with future

developments related to scientific results evaluation, could require reassessment and revision of the proposed CERIF extension/model in the future. Regarding the automation of the evaluation process, the constraint is that writing extensive and complex Jess rules could require the involvement of engineering experts. This could be resolved by applying new developments related to knowledge representation and reasoning over such representations.

Both constraints will be the main activity of our future work.

Acknowledgement

This work was supported by the Ministry of Science and Technological Development of the Republic of Serbia, through project no.III-47003.

References

- [1] A. Reinhardt and European Science Foundation, *Evaluation in Research and Research Funding Organisations: European Practices : a Report*, Strasbourg, European Science Foundation, 2012, ISBN: 9782918428831
- [2] Committee for the Evaluation of Research - CIVR, "Guidelines for Research Evaluation", Ministry of University and Research (MIUR) 2006, from: http://vtr2006.cineca.it/documenti/linee_guida_EN.pdf
- [3] M. Y. Vardi, "Revisiting the Publication Culture in Computing Research", *Communications of the ACM*, Vol. 53, No. 3, pp. 5-5, Mar. 2010
- [4] J. Russell and R. Rousseau, *Bibliometrics and Institutional Evaluation*, Encyclopedia of Life Support Systems (EOLSS) Part 19.3 Science and Technology Policy. Oxford, Eolss Publishers, 2002
- [5] J. Bar-Ilan, "Informetrics at the Beginning of the 21st Century—A Review", *Journal of Informetrics*, Vol. 2, No. 1, pp. 1-52, Jan. 2008
- [6] F. Moksony, R. Hegedűs, and M. Császár, "Rankings, Research Styles, and Publication Cultures: a Study of American Sociology Departments", *Scientometrics*, Vol. 101, No. 3, pp. 1715-1729, Dec. 2014
- [7] D. Surla, D. Ivanović, Z. Konjović, and M. Racković, "Rules for Evaluation of Scientific Results Published in Scientific Journals", *Management Information Systems*, Vol. 7, No. 3, pp. 003-010, 2012
- [8] R. Lister and I. Box, "A Citation Analysis of the ACE2005 - 2007 proceedings, with reference to the June 2007 CORE conference and journal rankings," in *Tenth Australasian Computing Education Conference (ACE 2008)*, Wollongong, NSW, Australia, 2008, Vol. 78, pp. 93-102
- [9] P. Küngas, S. Karus, S. Vakulenko, M. Dumas, C. Parra, and F. Casati, "Reverse-Engineering Conference Rankings: What does It Take to Make a Reputable Conference?", *Scientometrics*, Vol. 96, No. 2, pp. 651-665, Aug. 2013

- [10] K. Jeffery, N. Houssos, B. Jörg, and A. Asserson, "Research Information Management: the CERIF Approach", *International Journal of Metadata, Semantics and Ontologies*, Vol. 9, No. 1, pp. 5-14, 2014
- [11] V. Penca, S. Nikolic, D. Ivanovic, Z. Konjovic, and D. Surla, "SRU/W-based CRIS Systems Search Profile", *Program: Electronic Library and Information Systems*, vol. 48, No. 2, pp. 140-166, 2014, DOI: 10.1108/PROG-07-2012-0040
- [12] T. Seljak and A. Bošnjak, "Researchers' Bibliographies in COBISS.SI", *Information Services and Use*, Vol. 26, No. 4, pp. 303-308, Jan. 2006
- [13] S. Nikolić, V. Penca, D. Ivanović, D. Surla, and Z. Konjović, "CRIS Service for Journals and Journal Articles Evaluation", in *Proceedings of the 11th International Conference on Current Research Information Systems*, Prague, Czech Republic, 2012, pp. 323-332
- [14] A. Reinhardt and European Science Foundation, *Evaluation in Research and Research Funding Organisations: European Practices : a Report*. Strasbourg: European Science Foundation, 2012
- [15] S. Nikolić, V. Penca, and D. Ivanović, "Storing of Bibliometric Indicators in CERIF Data Model," presented at the ICIST 2013 - International Conference on Internet Society Technology, Kopaonik mountain resort, Republic of Serbia, 2013
- [16] D. Ivanović, D. Surla, and M. Racković, "A CERIF Data Model Extension for Evaluation and Quantitative Expression of Scientific Research Results", *Scientometrics*, Vol. 86, No. 1, pp. 155-172, 2011
- [17] A. Clements, B. Jörg, G. C. Lingjærde, T. Chudlarský, and L. Colledge, "The Application of the CERIF Data Format to Snowball Metrics", *Procedia Computer Science*, Vol. 33, pp. 297-300, 2014
- [18] D. A. Waterman, "How do Expert Systems Differ from Conventional Programs?", *Expert Systems*, Vol. 3, No. 1, pp. 16-19, Jan. 1986
- [19] J. W. Lloyd, "Practical Advantages of Declarative Programming", in *Joint Conference on Declarative Programming, GULP-PRODE*, 1994
- [20] D. Ivanovic, D. Surla, and M. Rackovic, "Journal Evaluation Based on Bibliometric Indicators and the CERIF Data Model", *Computer Science and Information Systems*, Vol. 9, No. 2, pp. 791-811, 2012
- [21] S. Nikolic, V. Penca, and D. Ivanovic, "System for Modelling Rulebooks for the Evaluation of Scientific-Research Results. Case Study: Serbian Rulebook", in *Proceedings of the 4th International Conference on Information Society and Technology (ICIST 2014)* Kopaonik, Serbia, 2014, pp. 102-107

An Information Security and Privacy Self-Assessment (ISPSA) Tool for Internet Users

Tomislav Galba¹, Kresimir Solic², Ivica Lukic¹

¹ J.J. Strossmayer University, Faculty of Electrical Engineering, Kneza Trpimira 2b, Osijek, Croatia, tomlav.galba@etfos.hr, ivica.lukic@etfos.hr

² J.J. Strossmayer University, Faculty of Medicine, Josipa Huttlera 4, Osijek, Croatia, kresimir.sollic@mefos.hr

Abstract: Privacy and overall information security are significantly affected by an Internet users' awareness, knowledge and behavior. Therefore, there is a user's awareness assessment needed before developing security solutions that will include the user of the information system. The present paper proposes a validated measurement instrument developed as a web based software security and privacy tool for self-assessment (ISPSA). This solution is based on a scientifically validated questionnaire, OWL ontology concept, evidential reasoning approach and the intelligent agent's algorithm. The main goal of this paper is to propose the solution that will raise awareness among Internet users on privacy and information security issues.

Keywords: awareness; Evidential Reasoning; Information Security; Intelligent Agent; OWL ontology; UISAQ; users

1 Introduction and Problem Statement

Many persistent and new problems regarding information security and user's privacy are causing the development of a wide range of new security concepts [1-4]. It is crucial for such new security solutions to consider the human component as well, due to the fact that users can significantly affect the overall security of an information system [5-7]. The security management should therefore integrate separate security areas in the overall security solution [8] combining risk management, infrastructure safety, hardware solutions, security protocols and users' education and control.

The most common approach to the user's low security awareness problem solution is education, as control could be considered unethical or hard to manage. The education of users on how to create better passwords [9], how to handle private data, how to be more careful towards unknown collocutors on the Internet [10],

should be based on some measurements of the users' current level of awareness and behavior. Several rare solutions for measuring users' awareness [11-13] and their potentially risky behavior [14, 15] were actually the reasons for building the present Information Security and Privacy Self-Assessment (ISPSA) tool for different users of the information systems, and more generally, for every user of the Internet.

The proposed solution is based on previously developed and validated information security questionnaire [16] combined into OWL ontology [17] with calculations relying on Evidential Reasoning approach [18] and back coupling, founded on the Intelligent Agent's algorithm [19].

This present paper is divided into following parts: Chapter two consists of a few subsections describing the complete background of the online self-assessment tool, referring to the validated UISAQ questionnaire, a brief introduction and the description of OWL ontologies, enhanced evidential reasoning algorithm and the description of the implementation of an intelligent agent. Chapter three provides a description of the complete solution with an output example and comments. Chapter four concludes the paper.

2 Background

2.1 Users' Information Security Awareness Questionnaire

The scientifically validated Users' Information Security Awareness Questionnaire (UISAQ) is a reliable measurement instrument [16]. The questionnaire has 33 items divided into two main scales: one scale measures the users' potentially risky behavior and the other measures the users' awareness. Each scale is divided into three basic subscales with 5 or 6 items presenting questions (Figure 1). For each question there are five answers proposed, graded on the Likert scale from one to five, where five means "good", as seen from the perspective of the information security and privacy engineer.

UISAQ possesses two more elements that do not belong to the validated part of the questionnaire. On the first page of the questionnaire there are demographic questions that can be changed regarding the research aims and the category of users, while on the last page there are two external questions and acknowledgments. Those two UISAQ elements were not needed and are not used as constructive elements in building the proposed self-assessment tool.

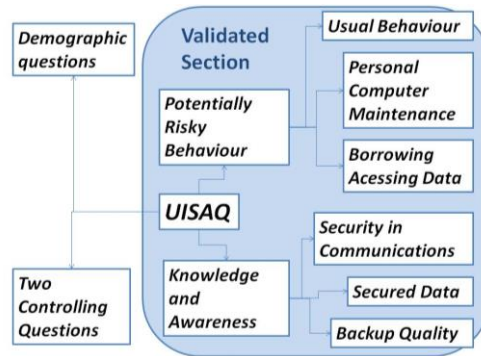


Figure 1
Segments of UISAQ

2.2 OWL Ontology

Ontology is used for formal definition of knowledge on some domain of interest. Formally, the defined knowledge should be both readable to a programmed intelligent agent and understandable to the human expert. In order to meet those requirements, there should be a language with well-defined semantics used [20]. There are many reasons for creating ontology [21]. Some of the most important reasons are:

- Reusability of domain knowledge which helps to develop a large and detailed ontology allowing integration with other new ontologies.
- Sharing common understanding for specific information domain as one of the main goals in ontology development that allows the extraction and aggregation of information from different domains.

The process of building an ontology is simple and often based on defining concepts with their properties and defining relations between concepts [17]. Nowadays, the most frequently used version of the ontology is OWL, while there are many open source software solutions to choose from. OWL is an international encoding and exchanging standard with the purpose of enabling communication between computers. It is developed as an extension on two semantic web standards, the RDF (Resource Description Framework) and RDF schema. Both of those standards were endorsed by W3C. So, naturally, OWL is a valid RDF document and also a well-formed XML document which facilitates processing with the already available XML and RDF processing tools and API-s. The OWL is also characterized as a higher level of expression relative to RDF, shown in Fig. 2. This property is very important in the process of describing attributes of some enterprises, since OWL described information is considered knowledge instead of just a simple data set.

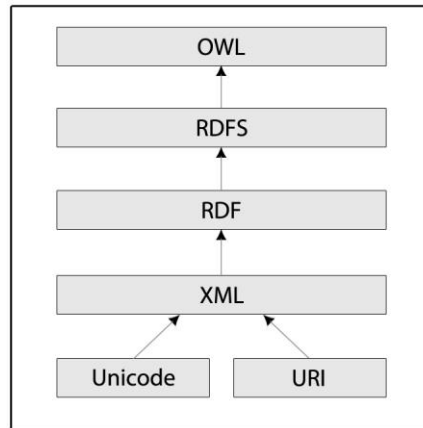


Figure 2
OWL hierarchy

There are three types of OWL expressive sub-languages [14, 22]:

- OWL Lite – is a subset of OWL DL that provides only the basics for subclass hierarchy construction which results in better performance of complete reasoners for OWL Lite.
- OWL DL – is a subset of OWL Full with less restrictions compared with the OWL Lite, designed to support description logic framework.
- OWL Full – very expressive language with no restrictions, designed as an extension to RDF. It contains all OWL language constructs with the possibility of unconstrained use of RFD constructs, resulting in full syntactic and semantic compatibility with RDF.

The ontology consists of classes, properties and individuals where [22, 23]:

- Classes are main building blocks representing sets of resources also known as individuals. Additional class description is achieved by adding properties from RDFS or OWL vocabularies. Also, a class can be related to other more general classes using `subClassOf` property which is shown in Figure 3.

```

<owl: Class rdf: ID="Class1" />
<owl: Class rdf: ID="Class2" />
  <rdfs: subClassOf rdf: resource="#Class1" />
</owl: Class>
  
```

Figure 3
Class definition

- Properties are used to define relations between individuals. There are two main categories of properties in OWL. The object property as an instance of predefined OWL class defines a link between individuals in two different classes, and Data type property that defines the relationship between the individual and data values. As with classes, properties can be described by adding sub elements such as subPropertyOf. Also, there are lots of other things we can add to the property, like establishment of taxonomy, domain and range of a property etc., shown in Figure 4.

```

<owl: ObjectProperty rdf: ID="hasPDescriptor" >
  <rdfs: domain rdf: resource="#P" />
  <rdfs: range rdf: resource="#PDescriptor" />
</owl: ObjectProperty >

```

Figure 4
Properties definition

- Instances – belong to classes and are used to express semantics of classes and properties. They are also related to other instances and data values by defining the properties. Two facts or axioms are used to define an instance:
 - Facts about properties of instances and class membership
 - Facts about identity of an instance

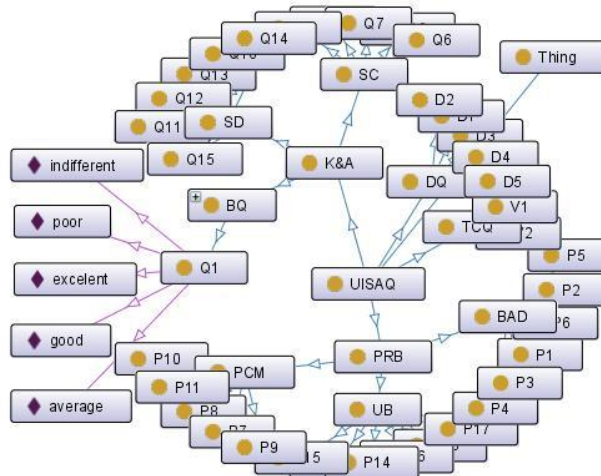


Figure 5
OWL ontology based on UISAQ built in Protégé

As stated before, classes are the main building blocks of ontology that contain instances (objects in the domain of interest). Classes are placed into superclass-subclass hierarchy. There are many benefits of using OWL ontology to formally define knowledge [17] such as re-usage among scientists and experts etc. Entities of basic subclasses are defined as grades from one to five, meaning from poor to excellent. The authors of this paper used the Protégé software solution [24] for building an OWL ontology, based on UISAQ questionnaire as shown in Figure 5.

2.3 Evidential Reasoning Algorithm

Evidential Reasoning Algorithm (ERA) is based on Dempster-Shafer theory [25, 26], decision-making theory [27] and evaluation analysis model [28]. It is also suitable for multiple-criteria decision analysis problems.

ERA is able to calculate with both quantitative and qualitative measurements considering subjective judgements with uncertainty, and objective, absolute or partial data [18]. ERA can be used to estimate the current state of technical system from many aspects. Estimated current state can be compared to the previous current state(s) of the same system, to current state(s) of another system or to previously defined referent value(s).

In this paper, the authors used the enhanced version of ERA [29]. Enhanced ERA allows aggregation of grades through a more complicated scheme with more than two levels regarding parent-child relation between the elements of the system.

The evaluation grades are defined in the same way as in the OWL ontology of UISAQ, namely as: poor, indifferent, average, good and excellent (P, I, A, G, E) while the uncertainty being calculated upon regarding of the missing answers.

The utility grade with associated utility interval is calculated from the distribution of grades with uncertainty and represents a single numerical value that is more suitable for comparing purposes. Utility interval is defined by uncertainty. Utility number can be calculated from any distribution of grades represented by any ontology subclass, single question or group of questions in the questionnaire.

An example of a calculation and aggregation through one ontology subgroup to a group of questions is shown in Table 1. More detailed explanation of ERA is available in wide scientific literature about this subject with many examples for technical systems state analysis [30-32], organizational decision making [33-35] and rarely for the evaluation of humans properties [14, 36, 37].

In order to apply ERA on ontology structure, three simple rules should be followed [14]: the hierarchical structure defined in the ontology should be strictly a cyclic graph; every direct relation should be “one-to-one” or “one-to-many” relationship; and there should be an existing crossing between classes in ontology reorganized. By additional reorganization it is possible to define mirrored classes in ontology.

Table 1
Example of calculations in the self-assessment process for the single user

UISAQ naming	Item	Subscale	Scale	Total	
Ontology naming	Class	Subclass	Superclass	Main class	
ERA naming	Basic attribute	Inter-attribute	Inter-attribute	General attribute	Utility (with uncertainty)
P1, P4	G	G(0.371)	I(0.085) A(0.045) G(0.222) E(0.618) H(0.031)	P(0.090)	U=0.772 (0.765-0.779)
P2, P3, P5	E	E(0.629)		I(0.038)	
P6	I	I(0.156)		A(0.114)	
P7,P8	G	A(0.156)		G(0.370)	
P14	E	G(0.344)		E(0.375)	
P16	E	E(0.344)		H(0.013)	
P17	A				
P9	-	I(0.138)			
P10,P12-P15	E	E(0.747)			
P11	I	H(0.115)			
Q1, Q3	E	P(0.191)	P(0.194) A(0.194) G(0.481) E(0.131)		
Q2	G	A(0.191)			
Q4	P	G(0.191)			
Q5	A	E(0.429)			
Q6,Q8	P	P(0.409)			
Q7,Q10	A	A(0.409)			
Q9	G	G(0.182)			
Q11-Q16	G	G(1.00)			

2.4 Intelligent Reflex Agent

When talking about intelligent agents and environments in which they act, we mainly perceive them as software or hardware implementations. Both of these implementations are based on some input from sensors which can also be software (user input form, data from database etc.) or hardware (temperature sensor, moisture sensor etc.), taking actions through actuators as a result of an analysis. An actuator can be a robotic arm, or in other cases actions can be taken on the software level either in terms of showing some data to the user, by changing the data in the database, or by an automatic creation of documents, reports etc. Mathematically speaking, the behavior of an agent is described with a function which transforms any input to adequate action [19]. There are four types of intelligent agents which satisfy the above-mentioned:

- Simple reflex agent – the simplest type of an intelligent agent where every action is based only on current input regardless of everything else.

- Model-based reflex agent – a more advanced type of an intelligent agent in relation to a simple reflex agent where action depends on current input from sensors and the history of previous actions for different inputs.
- Goal-based agent – an intelligent agent with predefined goals, similar to the model-based agent, only with the difference in checking the impact of a certain action on a defined goal.
- Utility based agent – operates through a utility function which is used to map a state. The result of that function is some kind of measure which defines how desirable a particular state of an agent is.

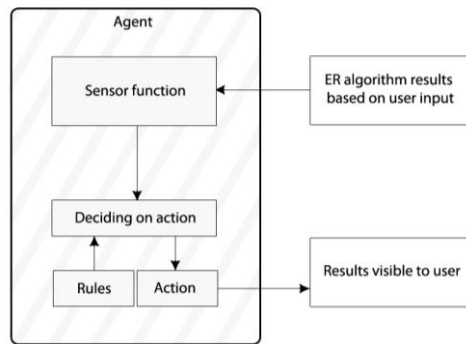


Figure 6
Simple reflex agent

This paper presents an ongoing work on the intelligent agent for online information security and privacy self-assessment tool. Since the main task of our intelligent agent in the self-assessment tool is to make decisions based on the final grade for the level of security, there is the model of a simple-reflex agent used. The use of this model is satisfying because there is no need to look at past, but only on current states. The expected output from the agent is information about whether to increase the level of security and the need to emphasize critical elements or sub-elements of the user behavior. The structure of our intelligent agent is shown in Figure 6.

Input variables for our intelligent agent are as follows:

- R_p – Referent value for average safe security level referring to the desired level of security which is predefined, based on previous testing and comparison with previous information system security assessment.
- U – Utility value given by information security assessment based on enhanced evidential reasoning algorithm.

Referent values are defined in previous testing conducted on the sample of 701 Internet users with different age, gender, technical knowledge, level of education, working position, coming from different institutions and business subjects [38]. The referent values are shown in Figure 8.

As stated in [19], a simple reflex agent brings simple decisions based on the current environment state, and since our intelligent agent is based on the simple reflex agent the decisions that had to be made are as follows:

- The proposed corrections of certain security segments if overall utility value is below the desired level of security.
- If the overall utility value is greater than the desired utility value including the correction value, then the intelligent agent searches for worst-rated basic attributes or items.

As an addition to the above stated, the intelligent agent has a predefined set of critical questions, to which, attention has to be paid, under all circumstances.

3 Software Web Solution

The present section introduces a new solution which implements all the elements stated earlier in this paper. One of the main goals of this work is to consider a person as negative influence on the information security system, in order to improve the current solutions and make future implementations better. Figure 7 shows a self-assessment tool structure with OWL Ontology structure based on UISAQ described in section 2.2 and the intelligent agent described in section 2.4.

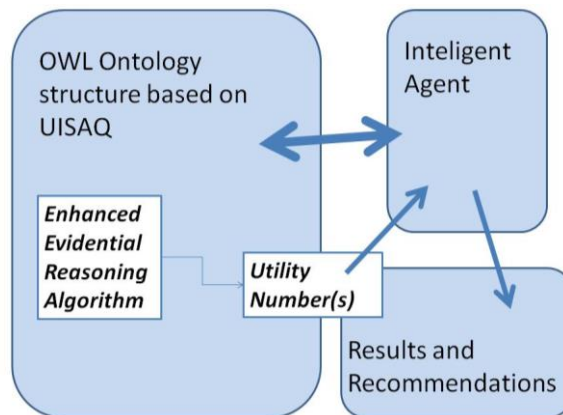


Figure 7

Self-Assessment Software Tool Elements

For a successful self-assessment, a user needs to pass over 33 predefined points divided into two major segments, which are then divided into six sub-segments, from which, every point has a different meaning (frequency, degree of security, degree of belief, degree of importance).

After passing through 33 points (it is not necessary to answer all questions), the algorithm for enhanced evidential reasoning is applied. The resultant values are shown in the form of a graph in Figure 8.

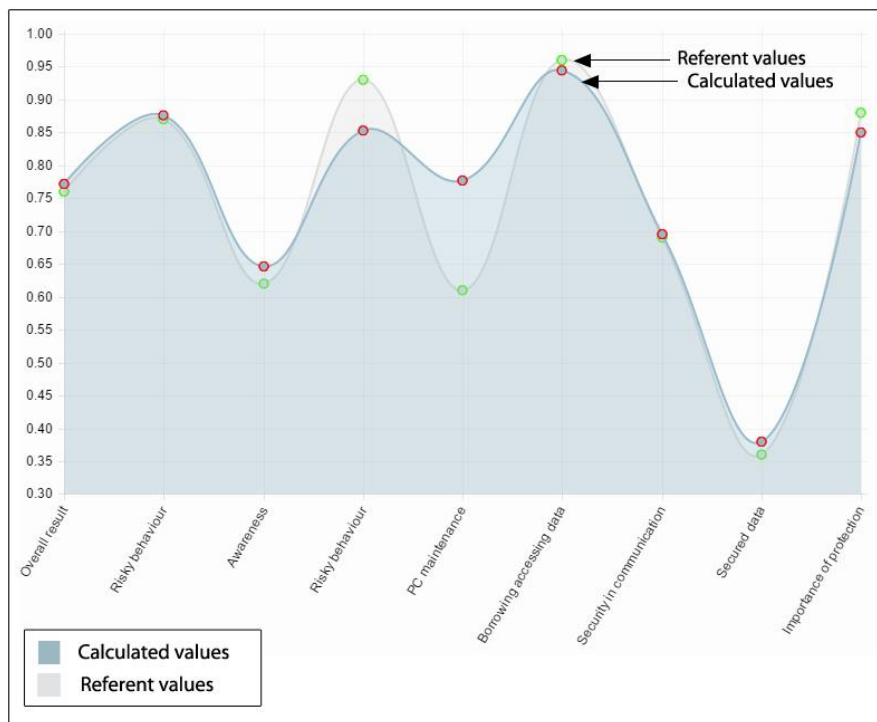


Figure 8
Self-assessment tool resulting graph

Figure 8 shows graph results for user input. From the calculated values we can conclude that the users have very good habits in most of the categories except few concerning the awareness, borrowing access data, protecting the data and poor backup habits. Except graph results, at the end, the user will the results from the intelligent agent, which emphasize the worst utility values and give recommendations for improvement. The working version of the questionnaire can be found in the referenced work¹.

¹ Towards Information Security and Privacy Self-Assessment (ISPSA) Tool for Internet Users link available at: <http://vns.etfos.hr/Samoprocjena/>

4 Discussion

The proposed self-assessment web solution tool has proven to be a valid measurement instrument that can be used to raise awareness among Internet users concerning privacy and information security issues. Once the user has passed through the 33 points in the self-assessment tool, the overall result is calculated. Also, the results for each of the two areas regarding behavior and knowledge, and the results for six subareas: usual behavior, PC maintenance, borrowing access data, security in communications, security of data and quality of backup are calculated. Moreover, there is an intelligent agent function which compares and analyses the overall user result with referent values representing the general user's behavior, knowledge and awareness called. Depending on comparison results, the user is pointed to critical security issues and provided with relevant recommendations for security improvement.

This Information Security and Privacy Self-Assessment Tool (ISPSA tool) for internet users is modular, based on scientifically validated UISAQ questionnaire, OWL ontology concept, enhanced Evidential Reasoning approach and intelligent agent's algorithm. ISPSA tool therefore benefits from each element's properties such as: measurement quality of the questionnaire, human machine utilization of the ontology, calculations with subjective assessment of evidential reasoning, and the agent's automation of analysis, as well as the presentation of results.

Future work will include some additional testing of the English version and making the self-assessment tool available freely to all Internet users. Also, with international collaboration, it should be possible to develop a better questionnaire, one which is more applicable to the world-wide Internet users' knowledge and habits and also more suitable to newly emerging information security issues. The modularity of the proposed solution also allows for the improvement of the different segments.

References

- [1] Sallai, G.: Future Internet Vision and Research Clusters, Acta Polytechnica Hungarica, 2014, Vol. 11, No. 7, pp. 5-24
- [2] Tot, L., Grubor, G., Takacs, M.: Introducing the Information Security Management System in Cloud Computing Environment, Acta Polytechnica Hungarica, 2015, Vol. 12, No. 3, pp. 147-166
- [3] Vokorokos, L., Baláz, A., Madoš, B.: Application Security through Sandbox Virtualization, Acta Polytechnica Hungarica, 2015, Vol. 12, No. 1, pp. 83-101
- [4] Vokorokos, L., Pekár, A., Norbert, A.: Yet Another Attempt in User Authentication, Acta Polytechnica Hungarica, 2013, Vol. 10, No. 3, pp. 37-50

- [5] H, Thompson: The Human Element of Information Security, IEEE Security&Privacy 2013, Vol. 11, pp. 32-35
- [6] Lukasik, S. J.: Protecting Users of the Cyber Commons, Communications of the ACM. 54, 9(2011) pp. 54-61
- [7] Solic, K., Sebo, D., Jovic, F., Ilakovac, V.: Possible Decrease of Spam in the Email Communication, IEEE MIPRO 2011, pp. 170-173
- [8] Michelberger Jr, P., Labodi, C.: After Information Security - Before a Paradigm Change (A Complex Enterprise Security Model), Acta polytechnica Hungarica, 2012, Vol. 9, No.4, pp. 101-116
- [9] Keszthelyi, A.: About Passwords, Acta Polytechnica Hungarica, 2013, Vol. 10, No. 6, pp. 99-118
- [10] I, Kirlappos., M. A, Sasse.: Security Education against Phishing: A Modest Proposal for a Major Rethink, IEEE Security&Privacy, 2012, Vol. 10, pp. 24-32
- [11] Vuković, M., Katušić, D., Skočir, P., Jevtić, D., Delonga, L., Trutin, D.: User Privacy Risk Calculator, Proceedings of the 22nd International Conference on Software, Telecommunications and Computer Networks, Split, 2014, pp. 1-6
- [12] Aggeliki, T., Spyros, K., Maria, K., Evangelos, K.: Process-Variance Models in Information Security Awareness Research, Information Management & Computer Security, 2008, Vol. 16, pp. 271 - 287
- [13] Petri, P., Mikko, S.: Improving Employees' Compliance through Information Systems Security Training: an Action Research Study, Puhakainen & Siponen Appendices, 2010, Vol. 34, pp. 757-778
- [14] Solic, K., Jovic, F., Blazevic, D.: An Approach to the Assessment of Potentially Risky Behavior of ICT systems' users, Tehnički vjesnik - Technical Gazette, 2013, pp. 335-342
- [15] Novakovic, L., McGill, T., Dixon, M.: Understanding User Behavior towards Passwords through Acceptance and Use Modelling, International Journal of Information Security and Privacy (IJISP) 2009, Vol. 3, pp. 9-27
- [16] Velki, T., Solic, K., Ocevcic, H.: Development of Users' Information Security Awareness Questionnaire (UISAQ) - Ongoing Work, MIPRO 2014
- [17] Horridge, M.: A Practical Guide To Building OWL Ontologies Using Protege 4 and CO-ODE Tools, The University of Manchester, 2011, URL: http://owl.cs.manchester.ac.uk/tutorials/protegeowltutorial/resources/ProtegeOWLTutorialP4_v1_3.pdf
- [18] Jian-Bo, Y., Dong-Ling, X.: On the Evidential Reasoning Algorithm for Multiple Attribute Decision Analysis under Uncertainty, IEEE Transactions on Systems, Man and Cybernetics, 2002, Vol. 32, pp. 289-304

- [19] Russell, S., Norvig, P.: Artificial Intelligence, A Modern Approach. Third Edition, PEARSON, 2010, pp. 46-50
- [20] Gali, A., Chen, C. X., Claypool, K. T., Uceda-Sosa, R.: From Ontology to Relational Databases. Shan Wang et al (Eds.): Conceptual Modeling for Advanced Application Domains, LNCS, 2005, Vol. 3289, pp. 278-289
- [21] Noy, N. F., McGuinness, D. L.: Ontology Development 101: A Guide to Creating Your First Ontology, Stanford Knowledge Systems Laboratory Technical Report KSL-01-05 and Stanford Medical Informatics Technical Report SML-2001-0880, 2001
- [22] OWL Web Ontology Language Overview available at: <http://www.w3.org/TR/owl-features/>
- [23] Korda, N., Astrova, I., Kalja, A.: Storing Owl Ontologies in SQL Relational Databases, Engineering and Technology, 2007, Vol. 29
- [24] Protégé software tool, Stanford Center for Biomedical Informatics Research. URL: <http://protege.stanford.edu/>
- [25] Dempster, A. P.: Upper and Lower Probabilities Induced by a Multivalued Mapping, Ann Math Stat, 1967, pp. 325-339
- [26] Shafer, G.: A Mathematical Theory of Evidence, Princeton University Press, 1976, New Jersey
- [27] Zhou, M., Liu, X. B., Yang, J. B.: Evidential Reasoning-Based Nonlinear Programming Model for MCDA under Fuzzy Weights and Utilities, International Journal of Intelligent Systems, 2010, pp. 31-58
- [28] Zhang, Z. J., Yang, J. B., Xu, D. L.: A Hierarchical Analysis Model for Multiobjective Decision making, Analysis, Design and Evaluation of Man–Machine Systems, Pergamon, Oxford, U.K, 1990, pp. 13-18
- [29] Yang, J. B., Singh, M. G.: An Evidential Reasoning Approach for Multiple Attribute Decision Making with Uncertainty, IEEE Trans Syst Man Cybern, 1994), pp. 1-18
- [30] Jagnjic, Z., Slavek, N., Blazevic, D.: Condition Based Maintenance of Power Distribution System, EUROSIM, 2004, pp. 13-14
- [31] Xin-Bao, L., Mi, Z., Jian-Bo, Y., Shan-Lin, Y.: Assessment of Strategic R&D Projects for Car Manufacturers Based on the Evidential Reasoning Approach, Int J Comput Intell Syst, 2008, Vol. 1, pp. 24-49
- [32] Zhang, X. D., Zhao, H., Wei, S. Z.: Research on Subjective and Objective Evidence Fusion Method in Oil Reserve Forecast, J Syst Simul, 2005, pp. 2537-2540
- [33] Beynon, M., Cosker, D., Marshall, D.: An Expert System for Multi-Criteria Decision Making Using Dempster–Shafer Theory, Expert Syst Appl, 2001, pp. 357-367

- [34] Wu, W. Z., Zhang, M., Li, H. Z., Mi, J. S.: Knowledge Reduction in Random Information Systems via Dempster–Shafer Theory of Evidence, *Inf Sci.* 174, 2005, pp. 143-164
- [35] Srivastava, R. P., Liu, L.: Applications of Belief Functions in Business Decisions: a Review, *Inf Syst Frontiers.* 5, 2003, pp. 359-378
- [36] Kong, G., Dong-Ling, X., Body, R., Jian-Bo, Y., Mackway-Jones, K., Carley, S.: A Belief Rule-based Decision Support System for Clinical Risk Assessment of Cardiac Chest Pain, *European Journal of Operational Research* 2012, pp. 564-573
- [37] Solic, K., Kramaric-Ratkovic, K., Antolovic, Z.: Applying Evidential Reasoning Approach in Biomedicine - Example on the APGAR Score, *Simpozij HDMI Medicinska Informatika*, Dubrovnik, 2013, pp. 7-9
- [38] Solic, K., Velki, T., Galba, T.: Empirical Study on ICT System's Users' Risky Behavior and Security Awareness, *IEEE MIPRO / ISS*, 2015, pp. 1623-1627

Fixture Design System with Automatic Generation and Modification of Complementary Elements for Modular Fixtures

Attila Rétfalvi

Subotica Tech, Marka Oreškovića 16, 24000 Subotica, Serbia; ratosz@vts.su.ac.rs

Abstract: Modular fixtures are usually built from standard elements that can be found in the modular element set of a certain manufacturer. But in some cases the user besides using these elements modifies some semi-finished elements that can also be found in the modular element set, or even the user produces some brand new elements. The motivation behind this can be to simplify the fixture, or to make the locating or clamping of a workpiece possible or more precise. This paper presents the methods of automatic generation, and if required, automatic modification of such new or semi-finished elements.

Keywords: Modular fixture; automatic element modification; CAFD (Computer Aided Fixture Design)

1 Introduction

Constructing an appropriate fixture that ensures the desired precision and stability is often a complicated and time-consuming task. That is the reason why many efforts have been made to develop a system that is able to automatically determine the number and the order of the setups for a given workpiece, and construct acceptable fixtures for each setup. Such a system could spare respectable time spent on fixture planning and design. The great majority of such attempts are based on the use of modular fixture elements. If the batch size is small, it is much quicker to build a modular than to produce a dedicated fixture. After the use the modular fixture can be dismantled, and stored in a considerably smaller place than a dedicated fixture. The elements of a modular fixture can be reused at a next fixture. But in some cases, if we use only the standard modular fixture elements found in a modular element set of a manufacturer, it would result in a very complicated or not sufficiently precise fixture.

In this paper the overview of the most known articles published in this field and the short introduction of the system developed by M. Stampfer and A. Rétfalvi will be presented. The system has been developed for box-shaped parts, first of all

cast gearbox housings not bigger than 1000x1000x600 mm. Box-shaped parts are often located with the help of one or two holes, so the problems that can occur during locating a workpiece over an inner cylindrical surface with the help of modular fixture elements will also be examined. In such cases when the problem cannot be acceptably solved with the help of the elements found in a modular element set, the users produce an extra (new) locator element, and complement the modular fixture built from modular fixture elements with that extra element. Thus, the fixture complexity can be considerably reduced or the locating precision can be increased. In modular fixture element sets there are some semi-finished plates (adapter plates) which are modified when some relationship tolerances justify their use. Thanks to this the number of needed setups or the precision requirements toward the fixture can often be reduced. In this paper, a method of automatic generation and modification of some of the above mentioned fixture elements will also be presented. The dimension, shape and relationship tolerances that can be achieved with the fixtures generated this way are in most cases in class IT7.

1.1 Literature Overview

Many researchers have been engaged with fixture design automation, some of them focused on feature recognition, and the conceptual solution of the fixture. Others focused on finding the optimal layout of the fixture elements. Some dealt with fixture verification, including stability, deformation and accessibility analysis. Some developed such a system that solves more of the above-mentioned problems. Bansal et al. [1], introduce an indirect feature recognition system that - starting from neutral STEP format of the workpiece model determines the removal volume, assuming that the stock is a boxlike hull of the model. It identifies slots, pockets and holes, and then tries to determine the best fixturing points on the workpiece (which ensure minimum tolerance deviations) by using a workpiece slicing technique at different user defined heights. Paris and Brissaud [2], present a process planning system that, after feature recognition assisted by an expert, associates machining processes to machining features and then organizes them into a global machining plan of the workpiece. Finally, it gives recommendations on fixturing features and determines the positioning quality, stability and cutter accessibility indices. Perremans [3], presents an expert system that builds a modular fixture if the fixturing features are given. In order to make the system manufacturer independent, he uses contact, assembly, and tightening features to describe the modular elements. Vichare et al. [4], introduce a model where the manufacturing resources like the machining tool or fixtures are represented in a unified manner for a CAx chain. Alacron et al. [5], developed a fixture planning and design system using the functional design theory. The user interactively prescribes the functional requirements, and the system, after defining the locating, supporting and clamping faces, selects modular elements and puts them on the defined place. Hou and Trappey [6], made a fixture design system, in their article the fixture layout reasoning and fixture element selection is described in detail.

Kumar et al. [7] introduce an interactive fixture design system in which the user defines the fixturing surfaces and the program builds an interference-free fixture. Mervyn et al. [8], show a fixture design system based on evolutionary (genetic) search algorithm. It is assumed that the imported workpiece model is well oriented according to tool axis, and the software tries to determine the number of required locating, supporting and clamping faces and points. Cecil [9], introduces an automated fixture design system, which groups the features having parallel directions (at planar surfaces their normal is taken into consideration while in case of holes their axis is of primary importance) in the same setup. After defining the number of setups, the system selects the locating and supporting surfaces, and selects the locating and supporting elements - and if it is required selects an ancillary supporting element - for each setup. Finally, the system checks if datum faces (faces connected with a relationship tolerance) are used for locating the workpiece, validates the clamping from the aspects of workpiece stability and from the aspect of collision between the tool and the workpiece. Hu and Rong [10], report a faster interference checking method between fixture elements, and between the fixture elements and the tool. The elements and the tool are substituted with simple geometric forms, and those forms are projected on a plane. Thus the problem is converted from 3D to 2D. Papastathis et al. [11] propose a fixturing solution, where not all clamps are static, some of them can change their position during machining, and this way increase the workpiece stability in case of thin-walled parts. Wan and Zhang [12], were also occupied with the problem of machining thin-walled parts. They propose that support layout should be determined through maximalization of the fundamental natural frequency of the workpiece – fixture system. An et al. [13], developed a dedicated fixture design system, which uses some standard elements. After the user defines the fixturing surfaces and points, the system selects some standard elements, and automatically adjusts the support height, then assembles all these on a base element. Jonsson and Ossbahr [14], present an overview on different kinds of reconfigurable and flexible fixturing solutions. Boyle et al. [15] give us a comprehensive review on different methods used for setup planning, fixture planning, unit design, and fixture verification. Wang et al. [16], give an overview on manufacturing fixtures and on automatic fixture design methods. Vasundra and Padmanaban [17], review the most recent works on machining fixture configuration planning, with special focus on the limitations of previous works in the determination of the elastic deformation of the workpiece, and the limitations in fixture layout optimization. Many different approaches to the fixture design and different stages of the fixture design have been introduced in these articles and reviews, and the overwhelming majority of the researchers used modular fixtures in their work. But none of them even mention the adapter plates and cases when the user supplements the modular element set with a locator element produced by the user himself. With the help of such supplementary elements and adapter plates we can often build a simpler fixture, and with the help of a special adapter plate, in some cases, we can even reduce the number of required setups.

2 Systematization of the Fixturing Tasks

This work is focused on the manufacturing of box-shaped parts, especially gearbox housings and their fixturing. The fixturing tasks, for box-shaped parts can be classified into three groups: supporting, locating and clamping.

2.1 Supporting Types at Box-shaped Parts

Gearbox housings are most often machined on horizontal machining centers. There are three supporting types that are most often used on horizontal machining centers (Figure 1):

- horizontal supporting (pos_1)*, where four sides of the workpiece can be machined in one setup
- vertical supporting (pos_2)*, where three sides of the workpiece can be machined in one setup
- special vertical supporting (pos_3)*, where three sides of the workpiece and the fourth side can partially be machined in one setup

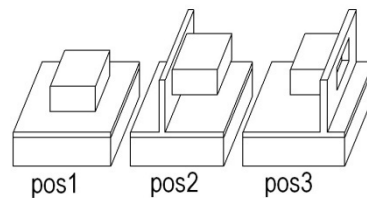


Figure 1
Types of the supporting [18]

2.2 Locating Types at Box-shaped Parts

There are four types of side locating established (Figure 2):

Type	Description	Sub-types			
p1	Locating by using surfaces which are on the adjoining faces of the plane locating face				
p2	Locating by using two inside diameter on the plane locating face				
p3	Locating by using one inside diameter on the plane locating face and a surface on one of the adjoining faces				
p4	Locating by using two threaded joints on the plane locating face				

Figure 2
Types of side locating [18]

- a) side locating with the help of surfaces adjoining the supporting face ($p1$),
- b) side locating with the use of two inside diameters on the supporting face ($p2$),
- c) side locating with utilization of one inside diameter laying on the supporting face and one face adjacent to the supporting face ($p3$),
- d) side locating with application of two threaded joints on the supporting face ($p4$).

2.3 Clamping Types for Box-shaped Parts

On the base of clamping force direction we can distinguish (Figure 3):

- a) perpendicular clamping ($s1$) – where the clamping force is perpendicular to the supporting surface
- b) parallel clamping ($s2$) – where the clamping force is parallel with the supporting surface
- c) clamping by screws and joints on the plane locating face ($s3$) - in this case the clamping forces are acting perpendicularly on the supporting surface, but the force transmission happens in a different way (in form closed manner).

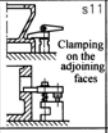
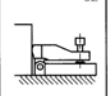
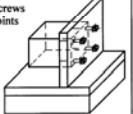
Type	Description of clamping type	Sub-types		
$s1$	Clamping perpendicular to the plane locating face	 $s11$	 $s12$	 $s13$
$s2$	Clamping parallel to the plane locating face	 $s2$		
$s3$	Clamping by screws and threaded joints on the plane locating face	 $s3$		

Figure 3
Types of clamping [18]

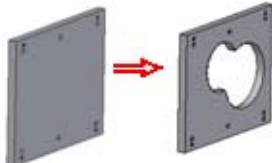
The basic type $s1$, depending on the location of the clamping faces, can be further divided into subtypes $s11$, $s12$ and $s13$. In the case of $s11$ the clamping surfaces are the closest parallel faces to the plane-locating (supporting) surface. In the case of $s12$ the clamping surface(s) is on the opposite side of the plane locating face. By $s13$ the clamping is carried out using a through hole on the workpiece.

The number of clamping points is also a very important characteristic of clamping. We distinguish clamping in one, two, three or four points. If we supplement the previous basic types with this information, we get that the possible clamping types are: $s11_2$, $s11_3$, $s11_4$; $s12_2$, $s12_3$, $s12_4$; $s13_1$, $s13_2$; $s2_1$, $s2_2$; $s3_2$, $s3_3$, $s3_4$. In the enumerated notation the last number denotes the number of clamping points.

3 Modular Element Sets

Process engineers noticed that some fixture elements occur in the same or in a slightly modified form in different fixtures. As it is much cheaper to produce similar elements in great quantities than producing somewhat different parts in piece production, engineers began to unify the most commonly used elements in order to be economical to manufacture them in great series. This led to the development of the modular element sets. The elements in such sets can be classified in three groups: base, functional and adapter elements.

Table 1
Adapter plates of a modular element set [20]

Element	Name
	Adopting plate (holes for bolts are machined according to the needs)
	Special plates

Base elements establish connection between machine tool table and the rest of the elements of the fixture. This group includes different palettes, grid plates and angled grid plates.

Functional elements are those elements that come in direct contact with the workpiece in order to fulfill a concrete task, such as supporting, locating or clamping; so this group includes different kinds of supports, locators and clamps.

Adapter elements are not used in every fixture, they are utilized when some supporting, locating or clamping surface is too high or too far and the gap between the functional element and the base element should be bridged. Adapter plates (Table 1) are used to make the locating or clamping of the workpiece possible or more simple.

3.1 Deficiencies of Modular Element Sets at Locating Workpieces over an Inner Cylindrical Surface

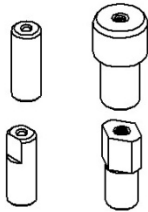


Figure 4

In modular element sets there are straight and flattened pins and bolts (Figure 4) used for positioning the workpieces on inner cylindrical surfaces. When one wants to locate a workpiece with the help of an inner cylindrical surface using modular elements, then he can encounter such difficulties, as there is no such locating element in the set whose diameter fits to the selected inner surface. For example, in the set of Heinrich Kipp Werk there are locating pins Ø8, Ø10, Ø12, Ø13, Ø14, Ø16, Ø18, Ø20, Ø22, Ø25, Ø30, Ø35, Ø40 and Ø50 mm. Thus, when a workpiece is to be located with the help of an inner cylindrical surface whose diameter is Ø28 mm one either has to produce a locator element (e.g. Ø27,95^{+0,02} mm) or has to use three smaller locating pins from the set to solve the problem. The inner cylindrical surface on the blank part is manufactured with a certain tolerance. As blank parts with lower tolerance limit (the smallest acceptable diameter) as well as blank parts with upper tolerance limit (the biggest acceptable diameter) have to be put on the fixture, the layout of the three pins must be such that matches the diameter of the lower tolerance limit. But if we put a blank part with upper tolerance limit on a fixture with such pin layout, the center of the locating surface can deviate from the ideal position. In Figure 4 the maximal locating error (Δx) can be seen in such cases. The blue circle (with radius R) presents the upper tolerance limit, the bold black line (with radius r) the lower tolerance limit, α is the half of the angle between two neighboring pins, and α_m is the half of the angle measured between the center of the blank part with upper tolerance limit and the centers of two neighboring pins, r_{cs} is the radius of the pins.

$$\Delta x = (R - r_{cs}) \cdot \cos(\alpha_m) - (r - r_{cs}) \cdot \cos(\alpha) \quad (1)$$

$$\text{where } \alpha_m = \arcsin\left(\frac{(r - r_{cs}) \cdot \sin(\alpha)}{R - r_{cs}}\right) \quad (2)$$

$$\text{and } \alpha = \frac{360^\circ}{2n} \quad (3)$$

(n - is the number of the pins used for defining the location of the workpiece)

In order to minimize the locating error Δx the three pins should be put on 120° to each other. It is very rare that every locating pin can be put in a grid hole of the base plate, usually at least one of them has to be moved at an appropriate place with the help of an adjustable support. As adjustable support elements should be screwed at least in two points to the base plate in order to avoid the rotation of these elements, the pins can be moved only along the directions shown in Figure 5, or possibly diagonally, if the adjustable support elements are long enough.

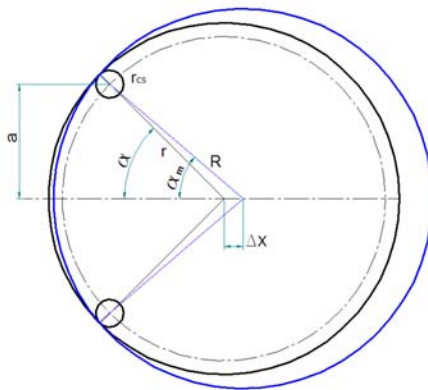


Figure 4
The change of the center point

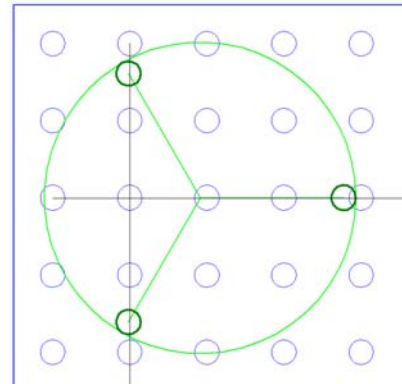


Figure 5
Adjusting locating pins

The big green circle represents the inner cylindrical locator surface, the three bold little circles illustrate the locating pins, and the little blue circles represent the grid holes on the base plate. With the use of adjustable support elements the precision and rigidity of the fixture are reduced. Due to grid hole step size (which is usually $50 \pm 0,01$ mm) it is often impossible to put the pins to be ideally 120° to each other; and at smaller diameters there might not be enough room for the adjustable supports. In certain cases there might even be surfaces on the workpiece that hinder putting the locator elements at the ideal place (the surfaces marked with blue color and letter **a**, and the missing parts of the **b** intermittent inner cylindrical surface). For example, if a gray cast iron workpiece made by gravity casting should be positioned over a $\varnothing 200$ mm inner cylindrical surface, and that surface after casting vary in diameter from 200 to 203 mm then the standard deviation of the position of the center can reach 3 mm if three locator pins are used to define the position of the centre. In this example $R=101,5$ mm and $r=100$ mm and we use $n=3$ pins with radius $r_{cs}=10$ mm, so $\alpha=60^\circ$, and α_m will be $58,41^\circ$, and finally Δx will be 2,93 mm. If a cylindrical locator plate $\varnothing 200$ mm is used for the same purpose Δx would not exceed 1,6 mm. The greater the Δx , the greater allowance on the blank part can be required in order to safely remove the so-called harmed layer during the roughing, or in order to ensure the minimal wall thickness after boring some holes. When the locating diameter (not covered by the element set) is smaller, it can easily happen that there is no other alternative than to produce a locator element since there is not enough space for pins and adjustable supports. There is another option, namely to sacrifice an adapter plate and machine holes in it for the pins thus avoiding the use of adjustable supports. However, the adapter plates, due to their size, are expensive and in many cases it is cheaper to lathe, mill, harden and grind an extra locator element.

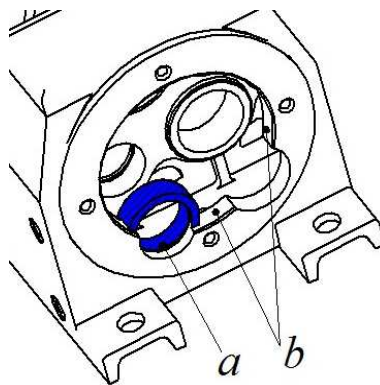


Figure 6

Example of hindering surfaces

From the above it can be concluded that, if in some cases the user complements the standard elements found in a modular element set with a locator element made by the user, it can result in a more precise, simpler, and often more rigid yet lighter fixture.

4 Setup and Fixture Design System

In order to speed up the fixture design process a Setup Planning and Fixture Design System (Figure 7) has been developed by the author and his collaborators. The input to the system is the CAD model of the workpiece saved in neutral IGES file format. The IPPO module with the help of Interface_1 can interpret the content of the textual IGES file, and thanks to the rules built in the CAD model post processor it recognizes the technological features on the workpiece model, and extracts the important geometrical information (like diameter, length, angle, etc.) for each of the features. The user interactively defines which surfaces, surface groups should be machined, gives the dimension, shape and relationship tolerances and the surface roughness to be achieved. The output of the IPPO module is the technological feature based model of the workpiece saved in a text file. The user opens that file with SUPFIX module, which with the help of Interface_2 can interpret the data in the opened file. The SUPFIX module verifies if all the surfaces to be machined can be finished in one setup on a horizontal machining centre. On a horizontal machining centre 4 sides of a workpiece can be machined in one setup, so box-shaped parts generally can be finished in one or in two setups. If the part cannot be finished in one setup, the program verifies if all relationship tolerance connected sides of the workpiece can be machined in the same (main) setup. If not, the program verifies if at least all strictly connected surfaces can be machined in the same (main) setup. If not, the strictly connected

surfaces have to be machined in different setups, but this requires more precise - therefore more expensive - fixtures. SUPFIX module gives recommendations on the supporting type and surfaces, on locating type and surfaces, and on clamping type and surfaces for each setup. All these recommendations – if accepted by the user – are saved in a text file. The user opens that file with the FIXCO module, which with the help of Interface_3 can interpret those saved data.

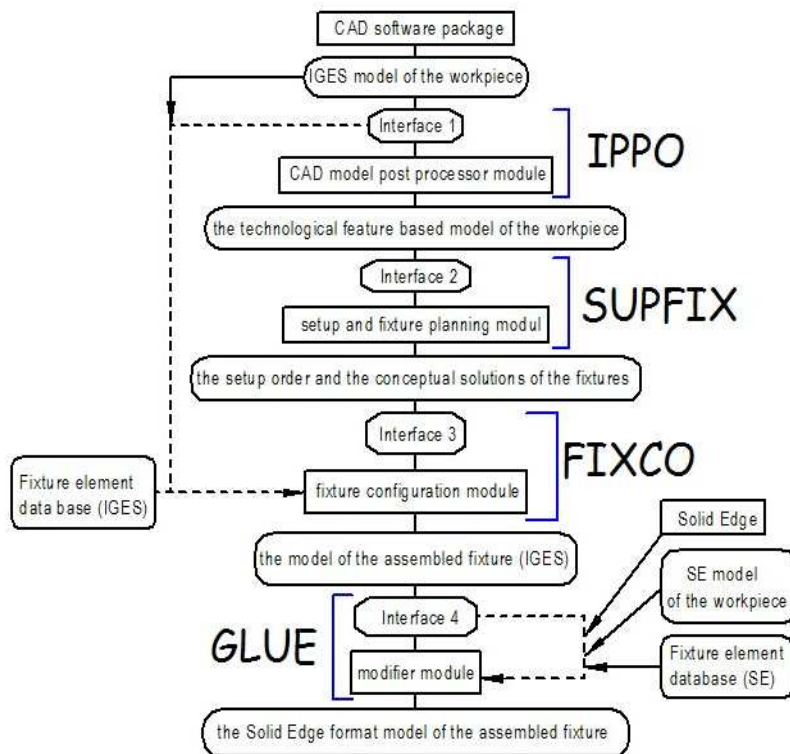


Figure 7
The Setup Planning and Fixture Design System

The FIXCO module builds a fixture for each setup by taking into consideration the recommendations given by the SUPFIX module. In the function of the proposed supporting, locating and clamping type and surfaces FIXCO selects and puts the needed fixture elements at the appropriate place. As FIXCO builds the fixture models from IGES format modular fixture elements, this module uses also Interface_1 to interpret and extract the needed data of the fixture elements. The final results of this activity are the fixture models for each setup built from IGES format modular fixture elements. The type, location, rotation of each fixture element – and, when necessary, the data required for generating or modifying a new non-standard element – are saved in a text file.

Finally this file is opened by the user with the GLUE module, which builds the CAD model of the fixture in Solid Edge Assembly Environment, and checks if there is interference between the fixture elements or between the fixture elements and the workpiece. If there is, it modifies the non-standard fixture elements to cancel the interference. The assemblies are built from standard modular fixture element models saved in Solid Edge's natural file format (.par files), but where it is necessary the GLUE generates the model of a new non-standard element, and where interference is to be avoided GLUE modifies the model of the non-standard elements.

The work of each module of the system is introduced in more detail in Stampfer [18], Rétfalvi [19], Rétfalvi and Stampfer [20]. In this paper, the automatic generation and modification of some complementary fixture elements is introduced in more detail.

4.1 Modification of the Fixture Elements

The elements that can be found in a modular element set do not always ensure the possibility of locating the given workpiece in a simple way. This problem most often occurs at positioning the workpiece on an inner cylindrical surface. In such cases the user has to produce an appropriate locator element. In other cases it is enough to modify some semi finished elements found in a modular element set. In modular element sets there are elements that are only partly machined, and their final shape must be made by the user, such elements being the simple and special adapter plates.

4.1.1 Generation and Modification of the Non-Standard Locator Elements

As mentioned, in modular element sets there are straight and flattened pins and bolts used for positioning the workpieces on inner cylindrical surfaces. But only certain diameters are covered with these elements, so if there is a hole with a diameter in-between two covered diameters, or one with bigger diameter than the greatest covered diameter and that hole would be the most appropriate to be used for positioning of a workpiece, sometimes the user must produce an extra (new) locator element. If the height (H_c) (Figure 8) of a such cylindrical locator element, its diameter (d) and the joining dimensions (M) are known the CAD model of that element can easily be generated automatically, and the element can also be put automatically at the appropriate place. The height H_c should be less than the height H_{max} (Figure 9) of the "upper" boundary curve of the locating surface, and must be at least 5 mm above the "lower" boundary curve of the locating surface ($H_c \geq H_{min}$). Diameter d of the locator element is equal to the lower tolerance limit (d_{min}) of the diameter of the locating surface.

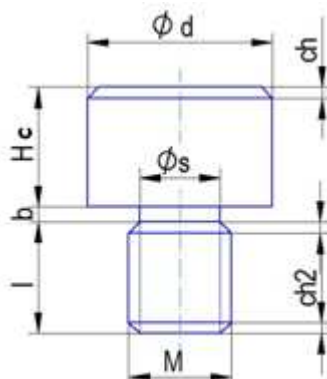


Figure 8
Cylindrical locator element

If the height of a locating element exceeded the double of the diameter, then to decrease the required height of the locator element, an adapter element (EA1, EA2, EA9 or EA10) should be used, and the locator element would be mounted in the adapter element. In this case, the height ($H_c = H_{min}$) of the locator element is of course measured from the upper surface of the adapter element. The generation of a (new) cylindrical locator element begins, when the proposed locating type is $p2$ or $p3$ (Figure 2) and the diameter of the proposed locating hole is not covered by the pins and bolts in the modular element set

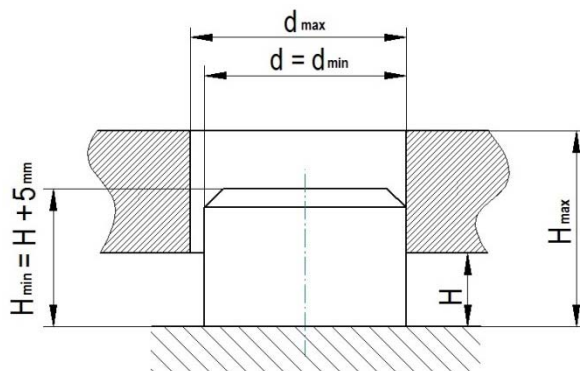


Figure 9
Height of the cylindrical locator element

The generation process of a cylindrical locator element (as shown in Figure 8) is: a circle with diameter d has to be extruded upward to height H_c , and the upper circular edge has to be chamfered with value ch . From the plane of the bottom surface another cylinder has to be extruded downward with diameter s and height

b , and from the bottom surface of the last extrusion another cylindrical extrusion is needed downward with diameter M on length l . Finally, the upper and lower circular edges of the last extrusion have to be chamfered with $ch2$. When the diameter of a non-standard locator element exceeds 75 mm, instead of downwardly extruding the cylindrical extension with $\varnothing s$ and $\varnothing M$, two counterbore holes (where the smaller diameter equals $\varnothing 12$ mm, and the larger diameter 20 mm) have to be made into it, to ensure the clamping and positioning of the locator plate to the base plate (Figure 10). The distance between the holes in z and y direction (where y and z are parallel to the supporting plane) must be dividable with the value of the grid spacing, and symmetric to the centre of the locator plate. A further complication can occur, especially at bigger diameters, for example, either a surface on the workpiece can hinder the correct locating (Figure 6) or some fixture elements cannot be placed on their ideal places because of another locator element (Figure 10). In such cases either the height of the locator element (H_c) is limited or one part of the locator element has to be removed. The limitation of the height means that the upper edge of the locator element is less than 5 mm above the lower boundary curve of the workpiece. It must be at least 1mm away from the hindering workpiece surface. In some cases the problem can be solved by reducing the height of the locator element so that it enters only 3 mm into the workpiece. However, before the height of the locator element is reduced, an approval from the user must be obtained. Practically this means that a cylindrical locator element with smaller height is generated – if approved by the user. The automated generation of cylindrical elements is relatively easy, but when the problem cannot be solved (due to some hindering workpiece feature or fixture element) by reducing the height of a locator element, it is a bit more complicated process to achieve the goal, namely to locate the workpiece over the proposed inner cylindrical surface. In such cases the program in the first step tries to find such a layout where there is neither interference between the workpiece and the fixture elements, nor between the fixture elements. If the program cannot find such a layout, it puts the models of all fixture elements and of the workpiece at an appropriate place (where every element can fulfil its task like supporting, locating or clamping) as if there were no interferences. The base concepts of the method are described in Rétfalvi and Stampfer [20]. The interference check verifying the elements which may have an inner point inside the cylindrical locator element is performed in the second step. The interference check is done by Solid Edge in Solid Edge Assembly environment, and is launched automatically by the GLUE module. Solid Edge saves in a textual file the list of the elements that are in interference with the locator element. In the third step the convex hull of the contour lines of the hindering workpiece or fixture elements are projected on the top flat surface of the locator element, and finally that projected forms are cut off with cut off distance H_c . The modification happens in automated way in Solid Edge Part environment, the .par file of the element to be modified is also opened automatically by the GLUE module.

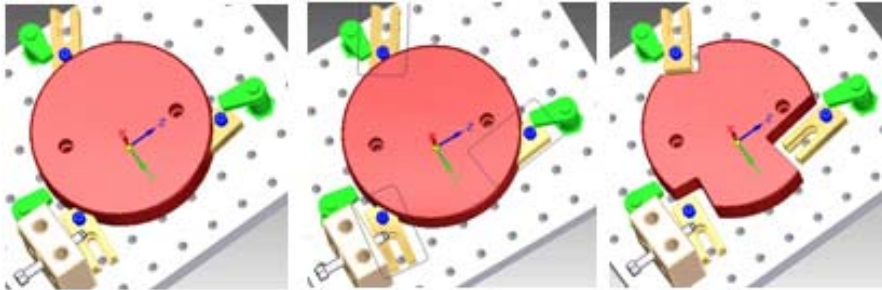


Figure 10

The process of modifying a non-standard locator element

Figure 10 shows an example of modifying a non-standard element; in the first picture the locator plate (generated in the first step) can be seen in interference with the supporting elements; in the second picture the 10 mm enlarged contours of the supporting elements are projected on the upper plane of the locator element; in the third picture the projected contours are cut off from the locator element. This process is activated if there is interference between a non-standard locator element and either any other fixture elements or the workpiece. This way (with extra locator element) a greater precision is achieved than with three pin locating. After cutting off some parts, of the locator element the clamps and supports can be put as close to each other as possible, thus the size and the weight of the whole fixture are smaller.

The flowchart of the automatic generation and modification of the extra locator elements is shown in Figure 11. If the type of the locating is neither $p2$ nor $p3$ (Figure 2) then it is either $p1$ or $p4$, so the elements used by the method described in this paper will not be used. If the locating type is $p2$ (and within it $p21$, $p22$, $p23$ or $p24$) then if there are pins or bolts with diameters that fit to inner cylindrical surfaces proposed (by SUPFIX module) for locating in the modular element set, and the height of these elements is in the range between $H_{\min}$ and $H_{\max}$ (Figure 9), the program selects two from these elements, and puts them to the appropriate places. If there are no appropriate locator elements, the program generates the missing ones with threaded ends (Figure 9). If the locating type is $p3$ the program examines if the locating can be solved with one standard locator element. It can be solved if there is a locator bolt or pin with the required diameter and height in the modular element set, and there are no such disturbing surfaces on the workpiece that hinder the use of the locator element. For example, in Figure 6, the bearing holder (a) hinders the use of such full cylinder locator elements whose height is greater than 4 mm (if the workpiece is to be located over the b intermittent inner cylindrical surface). Hence, if the locating type is $p3$, the program in the first step examines if there is such an element in the set whose diameter and height fit to the diameter and height requirements of the locating surface. If such an element is found the program selects and puts it to the

appropriate place. If there is no such an element, the program checks if there are three free grid holes inside the boundary curve of the locating surface. If this search is successful, the program tries to solve the locating with three pins and checks the positioning error (Δx). If Δx is greater than an acceptable value, instead of solving the locating task with three standard pins, the program tries to solve the problem with an extra locator element generated by the program in accordance with the dimensions of the proposed locating surface. The final decision whether the proposed solution is acceptable is made by the user.

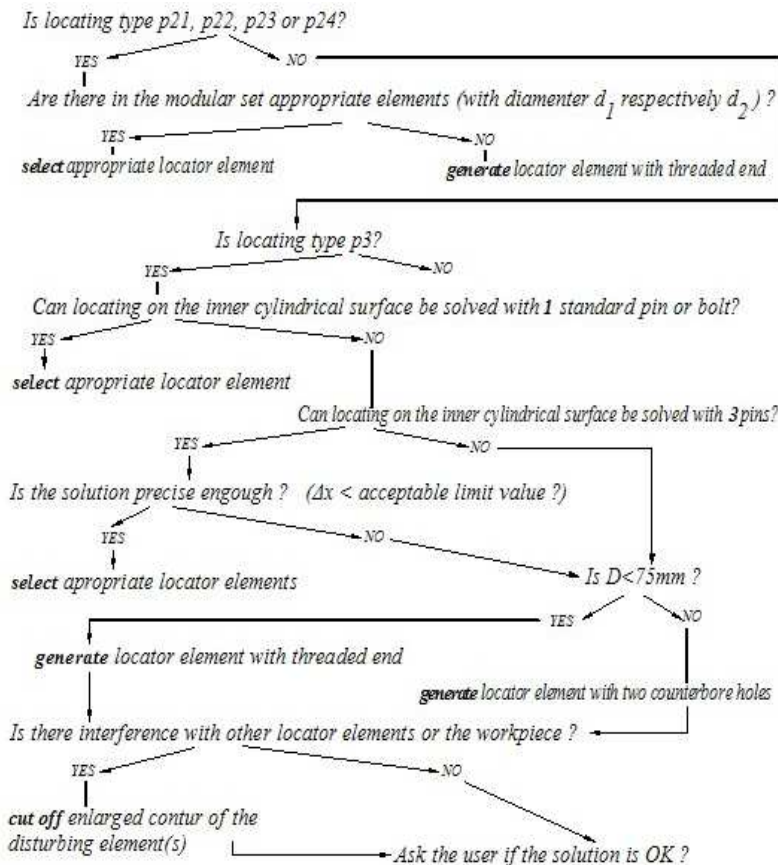


Figure 11

The outline of solving the locating problem

In Figure 13 on the left the final fixture (from the Figure 10) can be seen with the workpiece. Indeed the first proposal of the SUPFIX module for the auxiliary setup of this workpiece is to locate the workpiece over the small pink inner cylindrical surface (Figure 13, right), and if the user accepts that proposal, the program builds the fixture shown on the left picture in Figure 13. As there is no locator element in

the modular element set with appropriate diameter and height, the program generates an extra locator element with threaded end. At this case there is no need to cut off any part of the extra locator element. If the user refuses that first proposal, then the conceptual solution shown on the second small picture in Figure 13 is offered. Now the workpiece is to be located over the big red inner cylindrical surface (Pc1). As the locating error (Δx) is too big when the workpiece is located with the help of three bolts, an extra locator element is generated. Since the interference can in no way be avoided, some parts of the extra locator element are cut off automatically (see Figure 10).

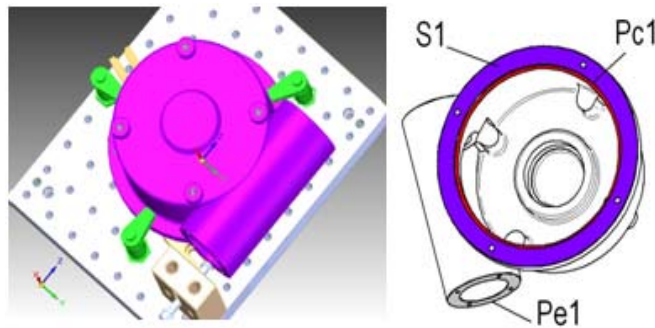


Figure 12

The fixture for the first setup from the Figure 11 together with the workpiece, and the adherent conceptual solution

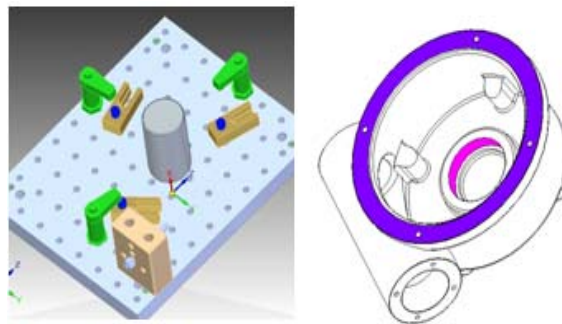


Figure 13

The first recommendation fixture for the first setup, and the adherent conceptual solution

4.1.2 Modification of the Adapter Plates

The modification of the adapter plates most of the time means making several holes on the appropriate places for some locator or clamping elements – for such purpose simple adapter plates are used. In other cases, besides these holes, some

openings have to be made, to enable the machining of some tolerance related features that lie on the supporting side of the workpiece – for such purpose special adapter plates are used. When the supporting type is *pos_1* or *pos_2*, simple adapter plates are used – if required; when the supporting type is *pos_3* then special adapter plates are used. Since FIXCO builds the fixture using IGES format workpiece and fixture element models, and these are hard to modify, the program at first makes the assembled model of the fixture with unchanged adapter plate, but stores the connecting dimensions and the locations of the elements (which demands connecting holes to be made). Then, the SE GLUE module is started which opens Solid Edge Assembly environment, and builds the Solid Edge model of the fixture assembly. If an adapter plate has to be placed in the assembly the program opens the CAD model of the adapter plate saved in Solid Edge's own *.par* format, and cuts the holes with appropriate dimensions on appropriate places, then saves the modified element under the same name, but with an end extension *_MOD* added, and just then puts the modified adapter plate to its right position in the CAD model of the fixture. This way the original adapter plate model can be used and modified later when building another fixture. For the automatic generation of the openings the somewhat increased contours of the surfaces, which lie on the supporting side of the workpiece, and should be machined at that setup, have to be projected on the supporting surface of the adapter plate. So if the supporting type is *pos_3*, the data of the contour curves (Rétfalvi (2011)) also have to be forwarded to the SE GLUE module. These curves lie within the inner boundary curve of the supporting surface of the workpiece, so a somewhat enlarged inner boundary curve of the supporting surface is cut out from the adapter plate. Figure 14 illustrates the second setup for the worm-gearbox housing, it is with four screws clamped to the adapter plate, the holes for four screws are to be machined by the user. The position of the holes is automatically determined and cut out from the model of the adapter plate.

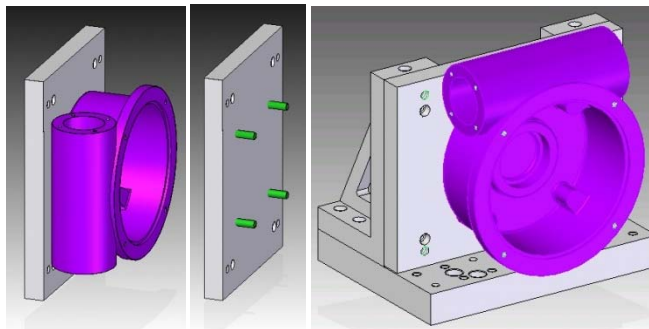


Figure 14

Adapter plate with and without the workpiece, and the entire assembled fixture for the second setup with the workpiece

Conclusions

Modular fixtures are built from standard elements found in a given modular element set, but sometimes are complemented with some extra elements, made by the user. In other cases so-called adapter plates, which are part of a modular element set, are modified by the user, and after the modifications are built into the fixture. In this paper, automatic generation and modification of the CAD model of these complementary elements and the automatic modification of the adapter plates is described. With the help of these complementary elements and the adapter plates the fixture configuration can often be simplified, and in some cases even the number of the required setups can be reduced. The method was developed for the setup and fixture design system also briefly introduced in this paper. With the help of this system a work that earlier took several hours can be completed approximately in half an hour's time. In Table 2 the time needed for problem solving (with AMD Athlon 64 processor, 3200+, 2,01 GHz and 2GB RAM) at different stages of fixture planning and design can be seen.

Table 2
The time needed for different stages of the fixture planning and design

Stages of the fixture planning and design	Gearbox housings							
	Part A		Part B		Part C		Part D	
Regeneration of the model and feature recognition	1 min 8 s		9 s		10 s		7 s	
Defining machining requirements	12 min		9 min		13 min		7 min	
Finding the conceptual solution of the fixture	1 min 42 s		1 min 8 s		2 min 10 s		1 min 56 s	
Fixture configuration	Auxiliary setup 56 s	Main setup 32 s	Auxiliary setup 46 s	Main setup 22 s	Auxiliary setup 41 s	Main setup 28 s	Auxiliary setup 30 s	Main setup 24 s

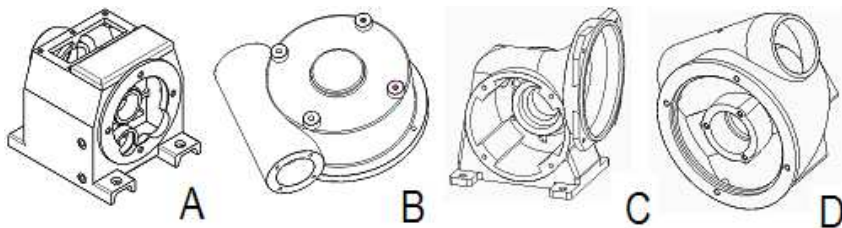


Figure 15
Gearbox housings

As the output of the system is the CAD model of the fixture with the clamped workpiece, the output file can be opened with a CAM module, and the fixture elements can be defined as check bodies before the toolpath generation. For fixture planning and design we must define which surfaces have to be machined with which precision, and for toolpath generation we also have to define which surfaces have to be machined. Taking this into consideration, the next step in the system development could be to ensure a greater level of integration with a CAM module in order to further reduce the technology planning time by eliminating the double intake of the same data (e.g. the surfaces to be machined).

References

- [1] S. Bansal, S. Nagarajan, N. Venkata Reddy (2008) An Integrated Fixture Planning System for Minimum Tolerances, *International Journal of Advanced Manufacturing Technology* 38, 501-513
- [2] H. Paris, D. Brissaud (2005) Process Planning Strategy Based on Fixturing Indicator Evaluation, *International Journal of Advanced Manufacturing Technology* 25, 913-922
- [3] P. Perremans (1996) Feature-based Description of Modular Fixturing Elements: the Key to an Expert System for the Automatic Design of the Physical Fixture, *Advances in Engineering Software* 25, 19-27
- [4] P. Vichare, A. Nassehi, S. T. Newman (2011) Unified Representation of Fixtures: Clamping, Locating and Supporting Elements in CNC Manufacture, *International Journal of Production Research* 49 (16) 5017-5032
- [5] R. H. Alacron, J. R. Chueco, J. M. P. Garcia, A. V. Idiope (2010) Fixture Knowledge Model Development and Implementation Based on a Functional Design Approach, *Robotics and Computer Integrated Manufacturing* 26, 56-66
- [6] J.-L. Hou, A. J. C. Trappey (2001) Computer-aided Fixture Design System for Comprehensive Modular Fixtures, *International Journal of Production Research* 39 (16) 3703-3725
- [7] S. Kumar, J. Y. H. Fuh, T. S. Kow (2000) An Automated Design and Assembly of Interference Free Modular Fixture Setup, *Computer-aided Design* 32, 583-596
- [8] F. Mervyn, A. S. Kumar, A. Y. C. Nee (2005) Automated Synthesis of Modular Fixture Designs using an Evolutionary Search Algorithm, *International Journal of Production Research* 43 (23) 5047-5070
- [9] J. Cecil (2004) TAMIL: an Integrated Fixture Design System for Prismatic Parts, *International Journal of Computer Integrated Manufacturing* 17 (5) 421-431

- [10] W. Hu and Y. Rong (2000) A Fast Interference Checking Algorithm for Automated Fixture Design Verification, *International Journal of Advanced Manufacturing Technology* 16, 571-581
- [11] T. Papastathis, O. Bakker, S. Ratchev, A. Popov (2012) Design Methodology for Mechatronic Active Fixtures with Movable Clamps, *45th CIRP Conference on Manufacturing Systems 2012, Procedia CIRP* 3, 323-328
- [12] X.-J. Wan, Y. Zhang (2013) A Novel Approach to Fixture Layout Optimization on Maximizing Dynamic Machinability, *International Journal of Machine Tools & Manufacture* 70, 32-44
- [13] Z. An, S. Huang, Y. Rong, S. Jayaram (1999) Development of Automated Fixture Design Systems with Predefined Fixture Component Types: Part 1, Basic Design, *International Journal of Flexible Automat Integrated Manufacturing* 7 (3/4) 321-341
- [14] M. Jonsson, G. Ossbahr (2010) Aspects of Reconfigurable and Flexible Fixtures, *Production, Engineering, Research and Development* 4, 333-339
- [15] Boyle, Y. Rong, D. C. Brown (2011) A Review and Analysis of Current Computer-aided Fixture Design Approaches, *Robotics and Computer-Integrated Manufacturing* 27, 1-12
- [16] H. Wang, Y. K. Rong, H. Li, P. Shaun (2010) Computer-aided Fixture Design: Recent Research and Trends, *Computer-Aided Design* 42, 1085-1094
- [17] M. Vasundra, K. P. Padmanaban (2014) Recent Developments on Machining Fixture Layout Design, Analysis, and Optimization using Finite Element Method and Evolutionary Techniques, *International Journal of Advanced Manufacturing Technology* 70, 79-96
- [18] M. Stampfer (2009) Automated Setup and Fixture Planning System for Box-shaped Parts, *International Journal of Advanced Manufacturing Technology* 45, 540-552
- [19] Rétfalvi (2011) IGES-based CAD Model Post Processing Module of a Setup and Fixture Planning System for Box-shaped Parts, *IEEE 9th International Symposium on Intelligent Systems and Informatics (SISY 2011)* 247-255
- [20] Rétfalvi, M. Stampfer (2013) The Key Steps toward Automation of the Fixture Planning and Design, *Acta Polytechnica Hungarica* 10 (6), 77-98

Comparison of Two Digital Cameras based on Spectral Data Estimation Obtained with Two Methods

Dejana Javoršek*, Tim Jerman, Andrej Javoršek

Faculty of Natural Sciences and Engineering, Department of Textiles, Chair of Information and Graphic Arts Technology, University of Ljubljana
Snežniška 5, SI-1000 Ljubljana, Slovenia; dejana.javorsek@ntf.uni-lj.si,
tim.jerman@fe.uni-lj.si, andrej.javorsek@ntf.uni-lj.si

*Abstract: The aim of our research was to prove that a digital camera does not influence the quality of spectral reflectance estimation and that satisfactory results could be achieved with a low-cost camera instead of using more expensive and complex multispectral devices. For that purpose two digital cameras – Nikon D300 in Nikon D700 were compared by obtaining spectral data from RGB values of a digital camera. For calculation of spectral data two different methods were used – SpecSens method and ImaiBerns method. In the research, two different color charts – ColorChecker DC and ColorChecker SG, were used. Performance of each camera and reflectance estimation approach were evaluated based on RMSE and ΔE^*_{ab} . Results showed that in the case of the ImaiBerns method it could be concluded that the obtained reflectance spectra are independent from the used camera, as somewhat slightly better results were obtained with Nikon D700. In the case of SpecSens method, which is based on the determination of the spectral sensitivity of the camera, the choice of the camera had quite an impact on the results. These results are pretty unreliable due to the large color differences ΔE^*_{ab} , as this calculation takes into account the standard light (first if you want to calculate XYZ and second the LAB values) and standard colorimetric observer.*

Keywords: digital camera Nikon D300; digital camera Nikon D700; spectral reflectance estimation

1 Introduction

Digital cameras have become very accurate systems for identifying changes in color, for example on beef [1], in the field of phenology [2, 3], in pattern recognition [4], for the identification of colors in urban environments [5], in the field of culture, where the digitization or digital archiving is used for the needs of various cultural institutions [6] or in medicine for photography of human wounds [7].

The sensitivities of digital camera differ from CIE color matching functions, which describe the sensitivity of the human visual system. Digital camera could provide two metamERICALLY identical images while human observers could see those images differently [8]. It is known that a color match for all observers when changing illumination could be achieved only by matching spectral data that are completely independent of the characteristics of digital camera. Obtaining spectral data from digital camera RGB values could provide a new way of using digital camera as spectrophotometric tool, where spectral data enables later reproduction under different illuminants and observing conditions. That could be very important and useful for digital archives, network museums, e-commerce and telemedicine [9]. Even after almost 25 years since Glassner wrote about deriving a spectrum from an RGB values [10], this is still a very hot topic. Today, one method to solve this problem is to use a regular digital RGB camera and estimate Suitable RGB image into a spectral image by the Wiener estimation method [11]. This method was also used for spectral reflectance images obtained from a digital RGB image for estimation of melanin concentration, blood concentration, and oxygen saturation in human skin tissue [12].

Adequate results of obtained spectral data could be also achieved by combining two different shots of the same scene acquired using the digital RGB camera with and without a properly chosen absorption filter [13].

There have been a number of studies on determining camera spectral sensitivity. For defining camera spectral sensitivity, a monochromator or narrow-band filters for generating a series of monochromatic light, are usually used. Other methods that do not use a monochromator require both input images and corresponding scene spectral radiances [14-19]. Thomson and Westland introduced a novel method to estimate camera spectral sensitivities and white balance setting from images with sky regions [20]. In our research instead of a monochromator, a less expensive and more readily available tools - diffraction grating and spectroradiometer were used to determine spectral sensitivities of two commercial digital cameras [21].

In one study authors researched the influence of camera parameters, e.g. exposure, on spectral data reconstruction of prints [22]. They found out that with multiple exposures it is possible to capture high dynamic range images, because limited dynamic range is the factor that lowers the reconstruction performance of consumer level cameras.

The aim of our research was to prove that the digital camera does not or only slightly influences the quality of spectral reflectance estimation and that satisfying results could be achieved with a low-cost camera instead of using more expensive and complex multispectral devices. For that purpose two digital cameras – Nikon D300 and Nikon D700 were compared while obtaining spectral data from RGB values of a digital camera. For calculation of spectral data two different methods were used. The first method was performed using spectral responses of a digital

camera (SpecSens method), and second one was Imai and Berns (ImaiBerns) method that includes linearized RGB values [23]. In the research, two different color charts – ColorChecker DC and ColorChecker SG, were used. ColorChecker DC was used only in case of the ImaiBerns method, as training set data and for the linearization method. Performance of each camera and reflectance estimation approach were evaluated based on root-mean-square-error (RMSE) and color difference equation ΔE^*ab .

2 Experimental

2.1 Materials and Methods

Color charts ColorChecker DC (X-Rite) – training set with 237 color patches – and ColorChecker SG (X-Rite) – test set with 140 patches, were photographed with two digital cameras – Nikon D300 in Nikon D700. Charts were illuminated by two light sources with color temperature 3194 K at 45° angle and distance of 170 cm. Camera settings were as follows – aperture: f/1.4, ISO value: 200, metering mode: Matrix, captured image format: RAW. Raw images were converted to TIFF files using an open source program dcrw [24].

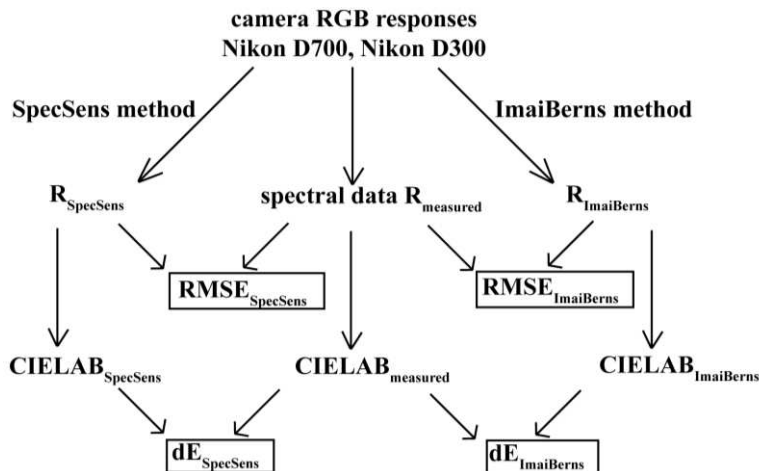


Figure 1

Workflow of the spectral reflectance estimation study, $dE = \Delta E^*ab$

Further characterisation steps, i.e. conversion from the RGB responses to the spectral reflectance data and other workflow details are depicted in Figure 1. Spectral reflectance estimation was carried out using two different approaches described below. The obtained results (denoted R_{SpecSens} and $R_{\text{ImaiBerns}}$) were compared to the actual reflectance values acquired by the spectrophotometer

EyeOne (denoted R_{measured}). Performance of the two reflectance estimation procedures was assessed using two error metrics – color difference formula CIE ΔE^*_{ab} and RMSE. For calculation of CIE ΔE^*_{ab} and CIELAB values illuminant A and CIE 1931 standard observer were used. RMSE is a spectral measure of estimation quality and compares measured and estimated reflectance values on image pixel location in each patch on a pixel-by-pixel basis.

2.2 Spectral Reflectance Estimation

The inverse problem of estimating spectral reflectances from the RGB values is related to the image acquisition process that describes the creation of camera responses. The process is described in [25]:

$$\mathbf{P} = \mathbf{f}(\Delta\lambda, \mathbf{Y}, \mathbf{l}, \mathbf{R}, \mathbf{b}) = \Delta\lambda * \mathbf{Y}' \times \text{diag}(\mathbf{l}) \times \mathbf{R} + \mathbf{b} \quad (1)$$

where $\Delta\lambda$ denotes the sampling interval of the spectral data, $\mathbf{Y}$ are the spectral responsivities of dimension $w \times c$, $\mathbf{l}$ is the illumination vector of dimension w , $\mathbf{R}$ is the reflectance matrix consisting of the n object pixels in the image and $\mathbf{b}$ is an additive noise term. In spectral reflectance estimation one attempts to calculate unknown reflectances $\mathbf{R}$ from known camera responses $\mathbf{P}$ by finding a function d that minimizes $d(g(\mathbf{P}), \mathbf{R})$. Here d is an error metrics, such as RMSE. Since dimensionality w of $\mathbf{R}$ is typically larger than c of $\mathbf{P}$, g does not necessarily have a unique solution [26].

2.2.1 Spectral Sensitivity-based Method (SpecSens)

For the purpose of our study, the Octave [27] function 'xyz2r' described in [23] that is used to estimate reflectance spectrum from tristimulus values was modified. Values for CIEXYZ and the standard illuminant/observer combination that are required as the function inputs were replaced by camera RGB responses and its sensors' spectral sensitivities.

In our research, to determine camera spectral sensitivities, a diffraction grating and a spectroradiometer were used instead of a monochromator [21]. A transmissive diffraction grating with 590 slits per mm was placed on the focal point of a biconvex lens onto which light from the illuminant A was projected. The diffraction grating split the parallel rays into several beams travelling in different directions depending on the wavelength. The resulting rainbow was photographed by the Nikon D300 and the Nikon D700 (Figure 2) and measured at several points by means of a spectroradiometer PR650 (X-Rite). From the obtained RAW data, the RGB values of the rainbow were read and a calculation of the corresponding wavelengths was performed. The interpolation of those points produced the spectral response curves for the camera. Since these curves still contained information about the input light spectrum, we divided the values by the interpolated measured intensity values from the spectroradiometer. Finally, the RGB curves were normalized so that the areas under the three curves were equal (Figure 3).

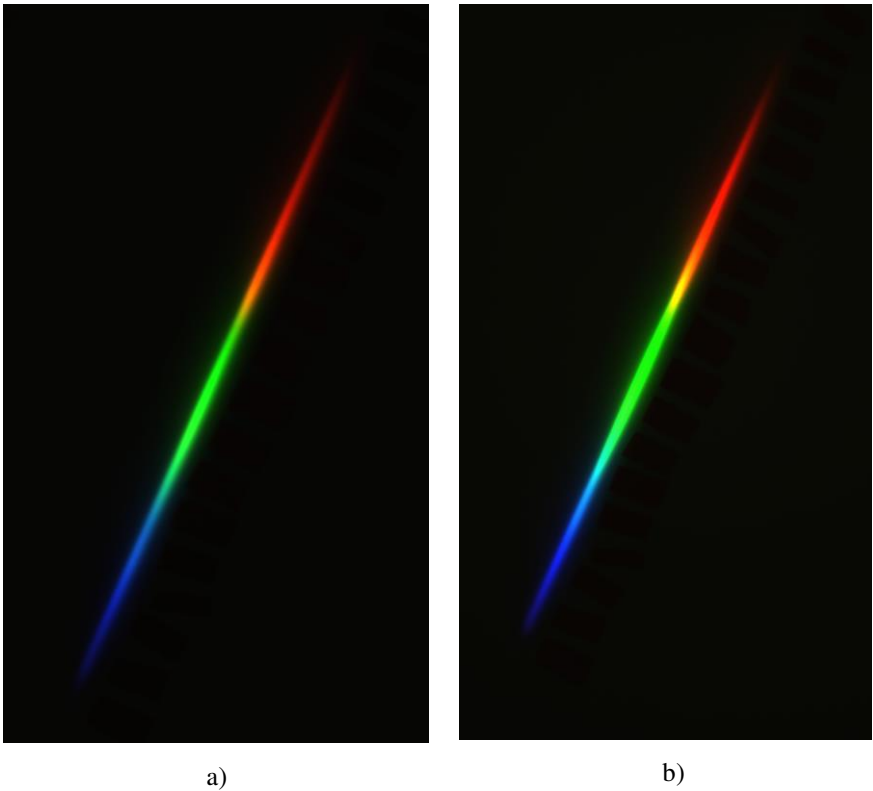
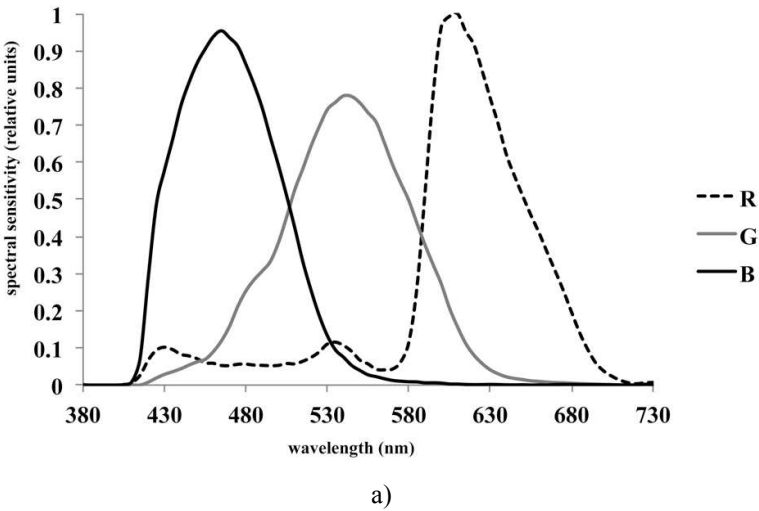


Figure 2
Rainbow photographed by a) Nikon D300 and b) Nikon D700



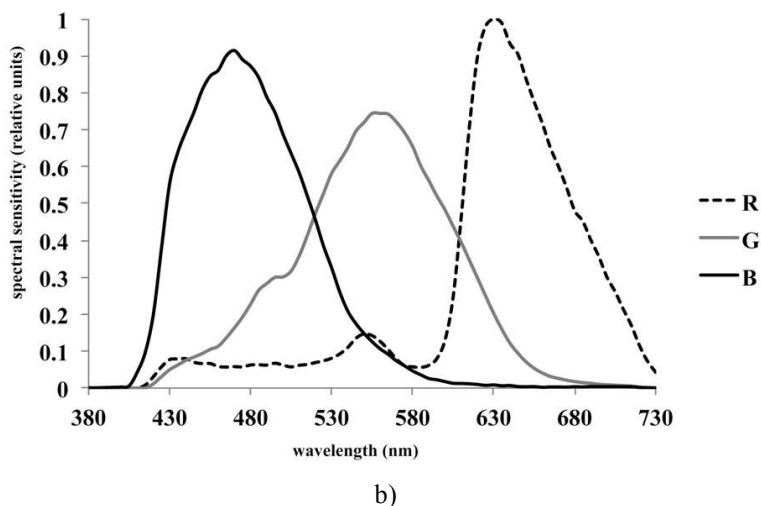
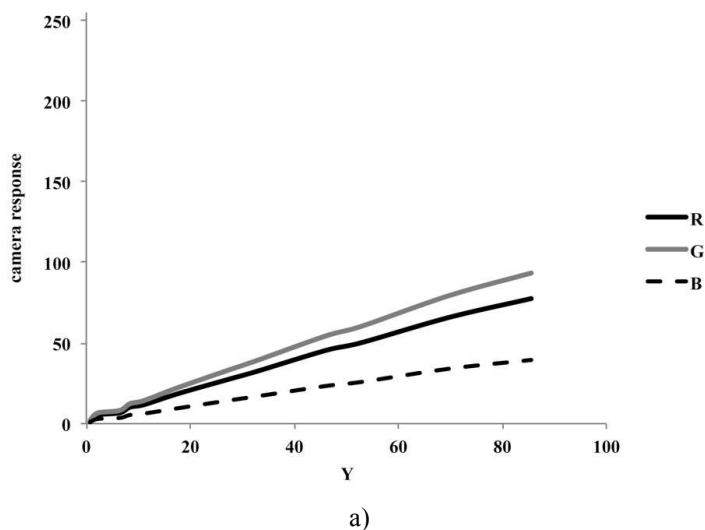


Figure 3

Spectral sensitivity curves of the camera sensors, a) Nikon D700 and b) Nikon D300

2.2.2 Imai-Berns Method (ImaiBerns)

First, linearization of the camera RGB responses was performed in order to compensate for the non-linearity introduced by the cameras manufacturer. The RGB values for 12 grey patches of the ColorChecker SG target and the corresponding CIE Y values representing lightness were used for this purpose (Figure 4). CIE Y values were calculated according to ISO 14524 [28]. Based on these data, linearization of all the 140 RGB values was accomplished using Octave functions 'getlincam' and 'lincam' [23].



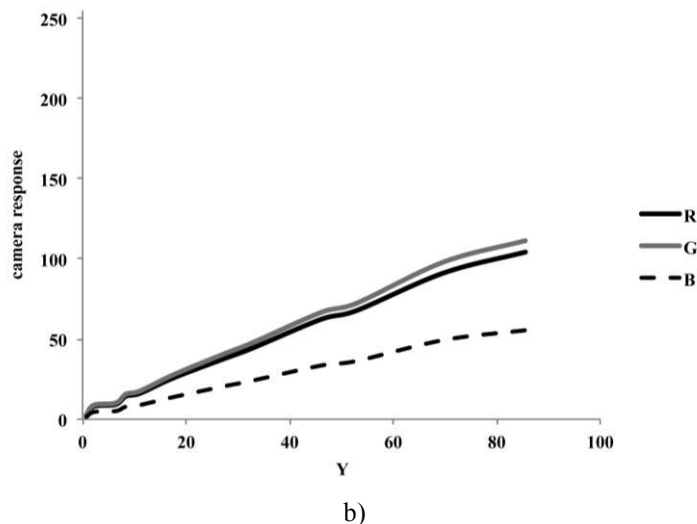


Figure 4

Linearization of the camera RGB responses, a) Nikon D700 and b) Nikon D300

Next, the linearized RGB responses were processed to obtain reflectance estimates using ImaiBerns approach [8, 21]. The method is based on running principal component analysis (PCA) on a training dataset. $\mathbf{R}_{tr}$ is expressed as a linear combination of k eigenvectors in a $w \times k$ matrix $\mathbf{V}$. The n_{tr} training reflectances are projected to the eigenvectors producing a $k \times n_{tr}$ coefficient matrix $\mathbf{E}_{tr} = (\mathbf{V})^T \cdot \mathbf{P}_{tr}$. $\mathbf{P}_{tr}$ is then related to the coefficient matrix $\mathbf{E}_{tr}$ as $\mathbf{G} = \mathbf{E}_{tr}(\mathbf{P}_{tr})^+$ using least squares regression. Superscript $+$ denotes the pseudoinversion of the $\mathbf{P}_{tr}$ matrix. Estimation of unknown reflectances from camera responses of the test set $\mathbf{P}_{te}$ is calculated as $\mathbf{R}_{tr} = \mathbf{V}\mathbf{G}\mathbf{P}_{te}$. The parameter k is obtained by exhaustive search with the goal to minimize the RMSE between the measured and the estimated reflectances in the training set of samples S_{tr} .

The same scene and illuminant were used for both digital cameras so mutual comparison could be achieved.

3 Results and Discussion

Results for two methods used for reflectance estimation in terms of two error metrics, ΔE^*_{ab} and RMSE, are represented in Table 1.

Comparison of both cameras using the colorimetric performance measure (ΔE^*_{ab}) showed better results for the Nikon D700 camera, but according to large values of ΔE^*_{ab} it is difficult to confirm that those results are reliable. In the case of the spectral metric RMSE, only the ImaiBerns method showed slightly better results

for Nikon D700 compared to Nikon D300. The reason for this could be found in the 12-megapixel FX (full-frame) sensor of the Nikon D700, which means it has a more dynamic range and higher ISO, while the Nikon D300 has a smaller 12-megapixel DX (1.5 crop factor) sensor.

Table 1
Performance of the methods used for reflectance estimation

	ΔE^*_{ab}		RMSE					
	ImaiBern	SpecSens	ImaiBern	SpecSens				
	D700	D300	D700	D300	D700	D300	D700	D300
mean	12.87	17.39	6.32	6.95	0.1332	0.1026	0.0596	0.0668
max	35.04	113.17	27.37	27.78	0.3856	0.2898	0.2150	0.2723
min	1.59	0.72	0.47	0.39	0.0132	0.0116	0.0015	0.0017

Both color difference calculations, ΔE^*_{ab} as well as spectral metric RMSE, clearly indicate that the reflectance estimation based on the linearized ImaiBerns method performs better compared to the alternative method (SpecSens) using defined camera spectral sensitivities. It should be noted that the differences between the two methods are almost the same when using both evaluation methods (RMSE and ΔE^*_{ab}): for example, in the case of ΔE^*_{ab} the corresponding ratio of mean values SpecSens/ImaiBerns is 2.04 (= 12.87/6.32) compared to 2.23 in case of the RMSE.

In order to get a more detailed picture about which color patches can be estimated more or less accurately in terms of their reflectance values, the following results are presented. Colors with the lowest ΔE^*_{ab} are colors with a mostly high L^* value. In order to get a more detailed information about color shift of these colors when their reflectance spectra were estimated from RGB values, they were represented in a^* , b^* diagram (Figures 5 and 6).

Color difference (ΔE^*_{ab}) between measured and predicted $L^*a^*b^*$ values using ImaiBerns method is the most pronounced in the case of black (patches L01, N06, N03, J05), saturated red (G04, M02, L03, L06, M05), blue (J04) and violet (M03) patch samples (Figure 5). Color difference (ΔE^*_{ab}) between measured and predicted $L^*a^*b^*$ values using SpecSens method is the most pronounced in case of black (patches I01, L01, N06, N03), saturated red (G04, M02, L03) and green (L07, L08, L09, G09, I09, H09) patch samples (Figure 6). Lightness of colors captured with Nikon D300 is slightly higher in comparison with the Nikon D700. Lightness of colors captured with the Nikon D700 and obtained with SpecSens method is slightly higher than in the case of the Nikon D300, which is quite opposite of the ImaiBerns method. However, as mentioned above, according to large values of ΔE^*_{ab} it is difficult to confirm that results of the SpecSens method are reliable.

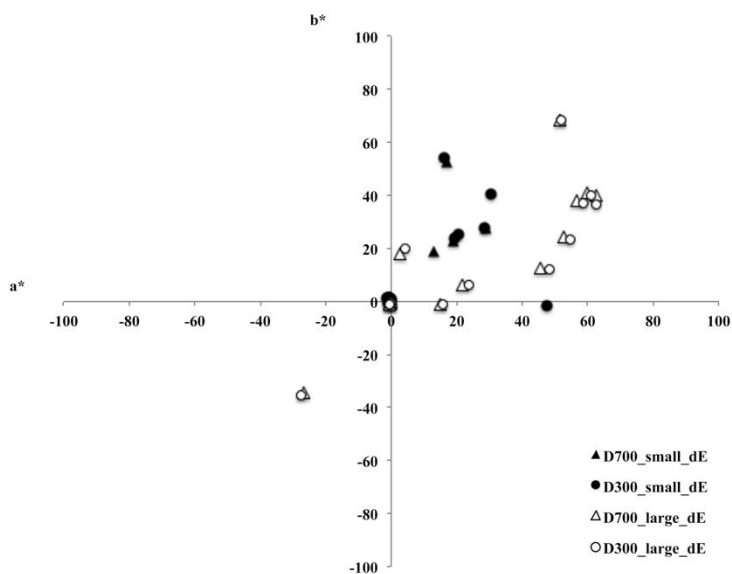


Figure 5

15 colors with the lowest and 15 colors with the highest ΔE^*_{ab} , obtained with the ImaiBerns method, represented in a^* , b^* diagram

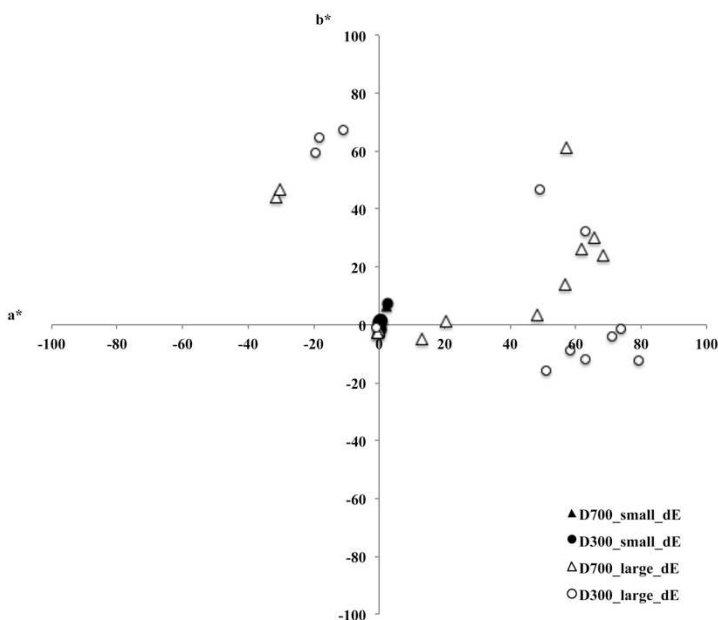


Figure 6

15 colors with the lowest and 15 colors with the highest ΔE^*_{ab} , obtained with the SpecSens method, represented in a^* , b^* diagram

On the other hand, it is also interesting to investigate patch samples with the lowest values of the ΔE^*_{ab} . Not surprisingly, patches with the lowest ΔE^*_{ab} are almost exclusively white, gray and skin color shade samples (ImaiBerns method).

Based on Figures 5 and 6 it can be concluded that in the case of the ImaiBerns method, the Nikon D700 colors with large ΔE^*_{ab} generally moved toward the b^* axis, meaning that they are less saturated than in the case of D300. In the case of the SpecSens method, large ΔE^*_{ab} colors of the digital camera D700 moved toward the upper half of the a^* , b^* diagram – their b^* values are positive, which means that colors moved from magenta area to red area of the diagram. In a case of yellow, colors moved from yellow area to green area of the diagram.

Figures 7-10 show the obtained spectral data of colors with the maximum and minimum ΔE^*_{ab} between the measured and the predicted $L^*a^*b^*$ values.

The spectral data represented in graphs confirmed previous findings that the difference between the cameras is small, except in case when SpecSens method was used, which unfortunately proved to be the less reliable method for comparing these types of cameras. From Figures 7 and 8 it is more clearly evident that small difference between the cameras is obtained in case of bright color (in Figure 7 white patch A04).

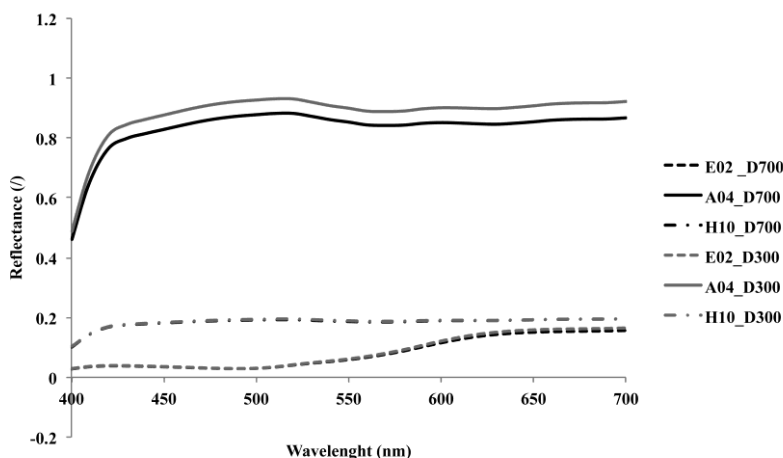


Figure 7

Reflectance spectra for six color patches obtained with the ImaiBerns method with small ΔE^*_{ab}

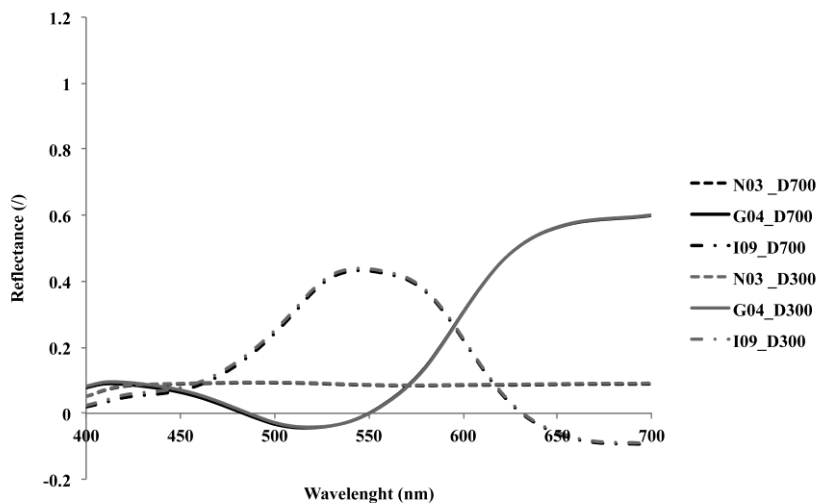


Figure 8

Reflectance spectra for six color patches obtained with the ImaiBerns method with large ΔE^*_{ab}

On the basis of spectral data in case of SpecSens method it could be seen that there is a difference between the cameras in both bright and dark colors, since for each camera separately spectral sensitivity was measured and calculated. This could lead to significant errors.

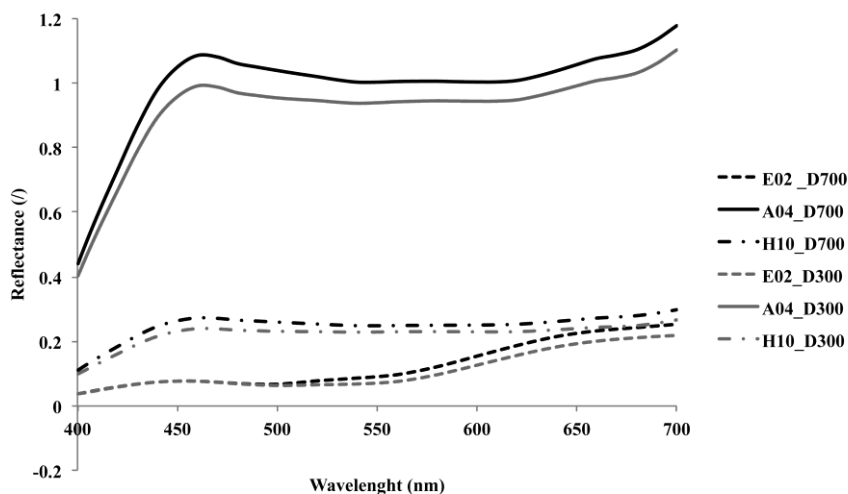


Figure 9

Reflectance spectra for six color patches obtained with the SpecSens method with small ΔE^*_{ab}

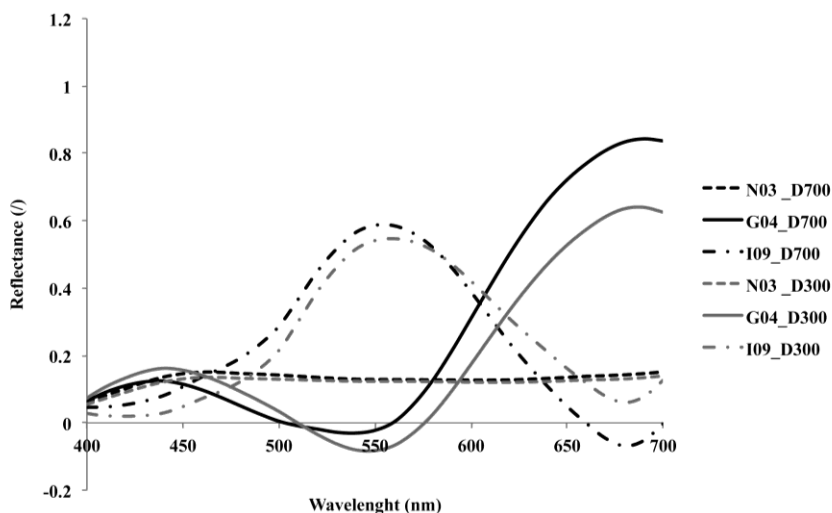


Figure 10

Reflectance spectra for six color patches obtained with the SpecSens method with large ΔE^*ab

Conclusion

In our study a comparison of two digital cameras Nikon D300 and Nikon D700 based on spectral data estimation obtained with two methods (ImaiBerns and SpecSens), was performed. Results showed that in the case of using SpecSens method and ΔE^*ab calculations, Nikon D700 had better results than Nikon D300. These results are pretty unreliable due to the large color differences ΔE^*ab , as this calculation takes into account the standard light (firstly when calculating XYZ and secondly when calculating LAB values) and standard colorimetric observer. In a case of SpecSens method measurement and calculation errors probably also occurred. In the case of the spectral metric RMSE, only the usage of the ImaiBerns method showed negligibly better results for Nikon D700 in comparison with Nikon D300.

Both evaluation methods, ΔE^*ab as well as RMSE, clearly indicate that the reflectance estimation, based on the linearized ImaiBerns method, performs better compared to the SpecSens method using defined camera spectral sensitivities. The reason for this could be found in the potential errors that occurred either in the actual measurements (SpecSens method) or calculations.

In the case of the ImaiBerns method it could be concluded that the obtained reflectance spectra are independent from the used camera, as somewhat negligibly better results were obtained with the Nikon D700.

According to ΔE^*ab calculations between measured and predicted $L^*a^*b^*$ values ImaiBerns method was the most pronounced in case of black, saturated red and blue patch samples, while SpecSens method was the most pronounced in case of

black, saturated red and green patch samples. Patches with the lowest ΔE^*_{ab} are almost exclusively white, grey and also skin color shade samples in the case of ImaiBerns method.

References

- [1] R. E. Larraín, D. M. Schaefer, J. D. Reed, "Use of Digital Images to Estimate CIE Colour Coordinates of Beef", *Food Research International*, Vol. 41, pp. 380-385, 2008
- [2] R. Ide, H. Oguma, "Use of Digital Cameras for Phenological Observations", *Ecological Informatics*, Vol. 5, No. 5, pp. 339-347, 2010
- [3] N. A. Clark, R. H. Wynne, D. L. Schmoldt, M. Winn, "An Assessment of the Utility of a Non-Metric Digital Camera for Measuring Standing Trees", *Computer and Electronics in Agriculture*, 28, No. 2, pp. 151-169, 2000
- [4] M. Tsutida, K. Yano, K. Hachimura, S. Tanaka, K. Furukawa, T. Nishiura, W. Choi, W. Wakit, H. T. Tanaka, "Development of a High-Definition and Multispectral Image Capturing System for Digital Archiving of Early Modern Tapestries of Kyoto Gion Festival", *Pattern Recognition (ICPR)*, 20th International Conference, Istanbul, Turkey, 2010
- [5] M. Starešinič, B. Simončič, S. Bračko, "Using a Digital Camera to Identify Colors in Urban Environments", *Journal of Imaging Science and Technology*, Vol. 55, No. 6, pp. 1-4, 2011
- [6] F. B. Wheeler, M. J. Bennett, "Accurate Color: a Preliminary Investigation into the Colour Gamut of Selected Special Collection Library Objects", *UConn Libraries Published Works*, accessed August 2011, available from http://digitalcommons.uconn.edu/libr_pubs/37
- [7] S. Van Poucke, "Automatic Colorimetric Calibration of Human Wounds", accessed June 2014, available from <http://www.biomedcentral.com/1471-2342/10/7>
- [8] F. H. Imai and R. S. Berns, "Spectral Estimation using Trichromatic Digital Cameras", in *Proc. of the International Symposium on Multispectral Imaging and Color Reproduction for Digital Archives*, Chiba University, Chiba, Japan, pp. 42-49, 1999
- [9] Y. Miyake, "Evaluation of Image Quality Based on Human Visual Characteristics", *Proc. of the First International Workshop on Image Media Quality and its Applications*, Nagoya, Japan, pp. 10-14, 2005
- [10] A. S. Glassner, "How to Derive a Spectrum from an RGB Triplet", *Computer Graphics and Applications*, IEEE, Vol. 9, No. 4, pp. 95-99, 1989
- [11] P. Stigell, K. Miyata, M. Hauta-Kasari, "Wiener Estimation Method in Estimating of Spectral Reflectance from RGB Images", *Pattern Recognition and Image Analysis*, Vol. 17, No. 2, pp. 233-242, 2007

-
- [12] I. Nishidate, T. Maeda, K. Niizeki, Y. Aizu, "Estimation of Melanin and Hemoglobin Using Spectral Reflectance Images Reconstructed from a Digital RGB Image by the Wiener Estimation Method", *Sensors*, Vol. 13, pp. 7902-7915, 2013
- [13] S. Bianco, F. Gasparini, R. Schettini, "Spectral-based Color Imaging using RGB Digital Still Cameras: Simulated Experiments", accessed June 2014, available from http://www.ivl.disco.unimib.it/papers2003/articolo_dsc_multispettrale.pdf
- [14] R. Slavuj, P. Green, "To Develop a Method of Estimating Spectral Reflectance from Camera RGB Values", *Colour and Visual Computing Symposium*, accessed June 2014, available from <http://www.google.si/url?sa=t&rct=j&q=&esrc=s&source=web&cd=34&ved=0CDsQFjADOB4&url=http%3A%2F%2Fcolorlab.no%2Fcontent%2Fdownload%2F42497%2F569619%2Ffile%2FCVCS-Radovan2013.pdf&ei=s6c1U-wDYeCzAPb8oHYDQ&usg=AFQjCNFbUTd62J68gtXT-9JSfwl23PB0Xg&bvm=bv.63808443,d.bGE>
- [15] P. M. Hubel, D. Sherman, J. E. Farrell, "A Comparison of Method of Sensor Spectral Sensitivity Estimation", *Proceedings of Color Science, System, and Application*, pp. 45-48, 1994
- [16] G. Sharma, H. J. Trussell, "Characterization of Scanner Sensitivity", *Proceedings of Transforms and Transportability of Color*, pp. 103-107, 1993
- [17] G. Finlayson, S. Hordley, P. Hubel, "Recovering Device Sensitivities with Quadratic Programming", *Proceedings of Color Science, System, and Application*, pp. 90-95, 1998
- [18] K. Barnard, B. Funt, "Camera Characterization for Color Research", *Color Research and Application*, Vol. 27, No. 3, pp. 153-164, 2002
- [19] M. Ebner, "Estimating the Spectral Sensitivity of a Digital Sensor using Calibration Targets", *Proceedings of Conference on Genetic and Evolutionary Computation*, pp. 642-649, 2007
- [20] R. Kawakami, H. Zhao, R. T. Tan, K. Ikeuchi, "Camera Spectral Sensitivity and White Balance Estimation from Sky Images", *International Journal of Computer Vision*, Vol. 105, No. 3, pp. 187-204, 2013
- [21] D. Javoršek, T. Jerman, B. Rat, A. Hladnik, "Assessing the Performance of a Spectral Reflectance Estimation Method based on a Diffraction Grating and a Spectroradiometer", *Coloration Technology*, Vol. 130, pp. 288-295, 2014
- [22] M. Nuutinen, P. Oittinen, "Recovering Spectral Data from Digital Prints with an RGB Camera using Multi-Exposure Method", *IS&T's Fourth*

- European Conference on Colour in Graphics, Imaging and Visualization 2010, pp. 120-125, 2010
- [23] S. Westland, C. Ripamonti, "Computational Color Science using MATLAB", Chichester: John Wiley & Sons, p. 181, 131, 134, 2004
- [24] "Decoding Raw Digital Photos in Linux", accessed July 2014, available from <http://www.cybercom.net/~dcoffin/dcraw/>
- [25] T. Eckhard, E. M. Valero, J. Hernández-Andrés, "A Comparative Analysis of Spectral Estimation Approaches Applied to Print Inspection", 18 Workshop of the German Color Group, Darmstadt, Germany, pp. 13-24, 2012
- [26] A. Ribes, F. Schmitt, "Linear Inverse Problems in Imaging", Signal Processing Magazine, IEEE, Vol. 25, No. 4, pp. 84-99, 2008
- [27] "GNU Octave", accessed July 2014 available from <http://www.gnu.org/software/octave/>; last accessed 8 July 2014
- [28] Photography – Electronic still-picture cameras – Methods for measuring opto-electronic conversion functions (OECFs) SIST ISO 14524:2011, Basel: ISO, p. 23, 179-185, 2011

Analytical Network Process in the Framework of SWOT Analysis for Strategic Decision Making (Case Study: Technical Faculty in Bor, University of Belgrade, Serbia)

**Živan Živković, Djordje Nikolić, Predrag Djordjević,
Ivan Mihajlović, Marija Savić**

University of Belgrade, Technical Faculty in Bor

Vojske Jugoslavije 12, 19210 Bor, Serbia

zzivkovic@tf.bor.ac.rs, djnikolic@tf.bor.ac.rs, pdjordjevic@tf.bor.ac.rs,

imihajlovic@tf.bor.ac.rs, msavic@tf.bor.ac.rs

Abstract: In this study Analytical Network Process (ANP) was applied as a model for prioritizing generated strategies based on the factors and sub-factors within the SWOT analysis, in the case of the Technical Faculty in Bor (TFB), University of Belgrade (UB), Serbia. ANP methodology approach implies the establishment of a hierarchical model on four levels: Goal (selection of the best strategy) - SWOT factors - SWOT sub-factors - alternative strategies, which establishes the interaction between clusters at different hierarchical levels of the model as well as the interactions between the elements within each cluster. This paper demonstrates a process for quantitative SWOT analysis that can be performed even when there is dependence among strategic factors. The proposed algorithm uses the ANP, which allows measurement of the dependencies among the strategic factors, as well as AHP, which is based on the independences between the factors. Dependencies among the SWOT factors and sub-factors are observed and their relative importance weights are determined, as well as their impact on the prioritization of the development strategy. The resulting benchmarking and prioritization of the alternative strategies in a series $WO_1 - SO_1 - ST_1 - WT_1$ for the development period of the TFB until 2025, indicates the sequence of application of certain strategies. This sequence implies that after reaching the limits in the application of the first strategy the next strategy in the defined sequence is implemented, in accordance with the mission of the TFB, the adopted strategic goals (SC) and adopted vision for the next ten-year period.

Keywords: ANP; SWOT; factors; sub-factors; strategy prioritization

1 Introduction

Strategic management includes a series of decisions and management actions in order to achieve the defined long-term goals of the company [1]. In this strategic management process a number of tools and techniques are used, among which analysis of Strengths, Weaknesses, Opportunities and Threats - (SWOT) has a special role [2]. SWOT analysis is a decision support tool and it is used as a tool for internal analysis as well as the analysis of the organizational environment. The obtained information can be systematically represented in a matrix, different combinations of the four factors from the matrix can aid in determining strategies for long-term progress [3-5]. However, this method does not provide analytical possibilities for quantification of the identified factors that are usually briefly and very generally described, and which represents a major disadvantage of SWOT analysis in strategic decision-making process [3, 6].

Through identifying strengths and weaknesses as a result of internal analysis, opportunities and threats as a result of the environment, organizations can build strategies that rely on strengths to reduce the perceived weaknesses, utilize identified opportunities and define a plan of actions to reduce or eliminate the impact of threats [7]. Acquired information can be systematically presented in the form of a SWOT matrix. In recent times different analytical methods have been developed which allow to determine the prioritization of long-term development strategies of the company, with different combinations of the four SWOT factors [3, 8, 9].

In order to eliminate weaknesses in the measurement and evaluation of steps in the SWOT analysis, a hybrid AHP method was developed [3], to quantify the weight of SWOT factors, which was named A'WOT in later studies [10-11]. This model has been tested in numerous studies, and despite its limitations, it is still widely used today [12-14]. AHP approach assumes that the factors presented in the hierarchical structure are independent. This approach can be questioned if the dependency between the SWOT factors is established, which can be determined by internal analysis of the organization and with the environment analysis [8].

The organization can make a good use of the opportunities if it has resources and strength to express its superiority, otherwise the opportunities will be lost because competitors will use them [8]. A similar relationship exists between strengths and threats. The ability to provide an adequate response to threats is based on the strength of the organization to eliminate or reduce the impact of threats. The dependency between the strengths and weaknesses in an organization is such that organizations with major strengths probably have fewer weaknesses, and therefore can more easily deal with situations which arise from the defined weaknesses. Organizations with more weakness relative to their competitors are more vulnerable to threats. These facts indicate that SWOT factors are not mutually independent, while factor weights of mutual connections depend on the specifics

of each organization [9]. Since the weight parameters are traditionally assign to factors as if there is no interdependence between them, under the conditions of existing interdependence the weight parameters can have different values, which directly affects the prioritization of the strategies [9, 15, 16].

Initial study developed in the eighties of the 20th Century [17] defined the AHP as a methodology of multicriteria decision making for complex problem solving. This method is a framework designed to cope with the intuitive, the rational, and the irrational when multi-objective, multi-criterion, and multi-actor decisions are made, with or without certainty for any number of alternatives. A basic premise of the AHP is the requirement of the functional independence of individual higher parts in the hierarchy from their lower parts, as well as between sub-factors within the same level. Many decision problems cannot be structured hierarchically because many factors from higher levels are related with lower levels as well as within the same level. Structuring problems with functional dependencies which allow feedback between the clusters represents Analytical Network Process (ANP). Saaty proposed the use of the AHP approach to solve problems in systems where there is independence between alternatives or criteria, and to use the ANP for systems in which there are direct or indirect connections between the individual levels [18]. For example, the importance of the criterion does not only affect the importance of the alternative, but also the importance of alternative affects the importance of the criterion. In addition, the elements of the cluster may affect some or all of the elements of any other cluster. Inner dependencies among the elements of a cluster are represented by looped arcs [9].

Implementation of the ANP methodology generally consists of four steps which are described in detail in the literature [8, 16]. ANP methodology is used in many complex systems in which interactions, in a hierarchical structure, occur between the level of clusters as well as between the elements within the cluster, which in the case of SWOT analysis clusters of factors and sub-factors could be used for prioritization of the strategies [8, 9, 15, 16].

In this study, ANP is used to define the relationship between the SWOT factors, SWOT factors and SWOT sub-factors, as well as between the sub-factors, for the purpose of the prioritization of the strategies. At the same time, the AHP method is used to determine the factor weights of dependency or independency and their influence on the selection of alternative strategies. The application of this methodology will be tested on a case of defining and prioritization of the strategies for the development of the Technical Faculty in Bor (TFB), University of Belgrade, Serbia by 2025. This tool, and obtained results, can be a starting point for benchmarking the operation of the TFB and further increasing its competitive position in the region.

2 Application of ANP Method on the Results of SWOT Analysis

The hierarchical and network model proposed for the SWOT analysis in this study consists of four levels, as shown in Figure 1. Goal (best strategy) represents the first level, the criteria (SWOT factors) represent the second level, sub-criteria (SWOT sub-factors) represents the third level and alternatives represents the fourth level (alternative strategies).

The hierarchical view of the SWOT model is shown in Figure 1a, while the general network model is presented in Figure 1b.

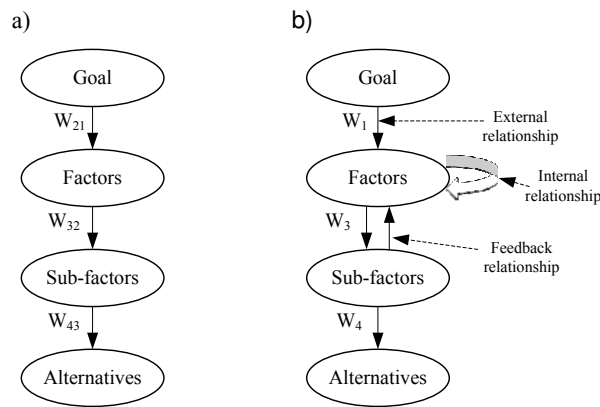


Figure 1

Comparison of the AHP and ANP structures - a) hierarchical model; b) network model

ANP model (Figure 1b) is an enhanced version of the AHP method, which more precisely defines the relationships in complex models that use many criteria, feedback and interdependence between the criteria. An advantage of this method is that it easily defines decision-making problem which includes many complicated relations. ANP method defines all components and relationships as bidirectional interactions. ANP includes relationships between individual clusters at different hierarchical levels, as well as the interactions between criteria and sub-criteria, therefore, this method is useful for obtaining more accurate and efficient results in decision-making in complex systems. Figure 1b illustrates that all criteria and clusters are interconnected through one of the potential links: unidirectional, bidirectional or loops. Unidirectional or bidirectional connection represents a connection between the clusters, while looping represents internal dependency in the cluster. The relative importance of the element i in relation to the element j is presented as:

$$a_{ij} = w_i/w_j \quad (1)$$

in the pairwise comparison matrix.

The pairwise comparison matrix A with n elements to be compared is formed as in eq. (3) [9]:

$$A = (a_{ij})_{n \times n} = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 1/a_{12} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ 1/a_{1n} & 1/a_{2n} & \dots & a_{nn} \end{bmatrix} = \begin{bmatrix} w_1/w_1 & w_1/w_2 & \dots & w_1/w_n \\ w_2/w_1 & w_2/w_2 & \dots & w_2/w_n \\ \dots & \dots & \dots & \dots \\ w_n/w_1 & w_n/w_2 & \dots & w_n/w_n \end{bmatrix} \quad (2)$$

After completion of the matrix A , assessment of the relative importance of the elements is performed by calculation according to the equation (4):

$$Aw = \lambda_{\max} \cdot w \quad (3)$$

where:

$\lambda_{\max}$ – largest eigenvalue of the matrix A ,

w – desired estimate.

AHP and ANP are popular methods also because they have the ability to identify and analyze inconsistencies of decision makers in the process of discernment and evaluation of the elements of the hierarchy [17]. If the values of the weight coefficients of all the elements that are mutually compared at a given hierarchy level could be precisely determined, the eigenvalues of the matrix A would be entirely consistent, however, this is relatively difficult to achieve in practice. Therefore, the application of these methods provides the ability to measure errors of judgment by computing consistency index (CI) for the obtained comparison matrix A , and then to calculate the consistency ratio (CR) [18].

In order to calculate the consistency ratio (CR), the consistency index (CI) needs to be calculated first, according to the following relation:

$$CI = \frac{\lambda_{\max} - n}{n - 1} \quad (4)$$

Then, the consistency ratio is determined by equation:

$$CR = \frac{CI}{RI} \quad (5)$$

where RI is a random index which depends on the order n of the matrix A , and is taken from the Table 1 [18].

Table 1
Random indices (RI)

n-order of the matrix A	1	2	3	4	5	6	7	8	9
RI	0.00	0.00	0.58	0.90	1.12	1.24	1.32	1.41	1.45

If the consistency ratio (CR) is less than 0.10, the result is sufficiently accurate and there is no need for adjustments in the comparisons and recalculation of the weights. However, if the consistency ratio is greater than 0.10, the results should be re-analyzed and the reasons for the inconsistencies should be identified and then removed by partial repetition of the pairwise comparison.

ANP approach consists of the following three matrices: supermatrix, weighted supermatrix and limit matrix. In the supermatrix the relative importance of all components is provided, in the weighted supermatrix the values obtained from the supermatrix of each cluster are defined. In the limit matrix, the constant values of each value are determined by taking the necessary limit of the weighted supermatrix [9]. The results of the decision making problem is obtained from the limit matrix scores [18].

The basic steps in the proposed SWOT-ANP model consist of the following. In the first step, SWOT factors, SWOT sub-factors and alternatives are identified. The procedure of obtaining importance of the SWOT factors, which represents the first step of the matrix manipulation concept of the ANP concept, is fully described in the literature [12, 19, 20, 21]. According to inner dependencies between the SWOT factors, inner dependency matrix is obtained and used to correct SWOT factor matrix. Then, SWOT sub-factors weights and priority vectors for alternative strategies are determined. Based on the schematic representation on Figure 1b the general supermatrix for the SWOT model which was used in this paper has the following form:

$$W = \begin{matrix} & \begin{matrix} \text{Goal} \\ \text{SWOT factors} \\ \text{SWOT sub-factors} \\ \text{Alternatives} \end{matrix} \end{matrix} = \begin{bmatrix} 0 & 0 & 0 & 0 \\ w_1 & W_2 & 0 & 0 \\ 0 & W_3 & 0 & 0 \\ 0 & 0 & W_4 & I \end{bmatrix} \quad (6)$$

where:

w_1 - vector that represents the impact of the goal on the selection of the best strategy based on the SWOT factors

W_2 - matrix that indicates the internal interdependence of the SWOT factors

W_3 - matrix which indicates the influence of the SWOT factors on the SWOT sub-factors,

W_4 - matrix that identifies the impact of the SWOT sub-factors on the alternatives.

It is preferable to present the details of the obtained results in this algorithm by using matrix operations.

In order to apply the ANP in the matrix operations for the purpose of determining the priorities of identified alternative strategies based on the SWOT analysis, the following steps are recommended [8, 9, 15]:

Step 1) Identify SWOT sub-factors and determine the alternative strategies according to SWOT sub- factors.

Step 2) In this step the importance of each SWOT group (strengths, weaknesses, opportunities and threats) is determined by calculating the weight matrix w_1 , while considering the situation that there is no internal interdependence between the SWOT factors.

Step 3) Calculation of the W_2 - inner dependence matrix of SWOT factors, using a scheme of internal interdependence shown in Figure 2

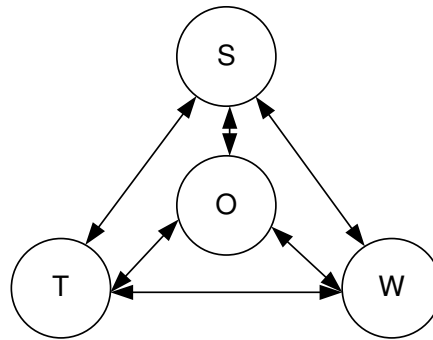


Figure 2
Internal interdependence of SWOT factors

Step 4) Calculating the weight matrix $w_{\text{SWOT factors}} = W_2 \times w_1$, of interdependent priorities of SWOT criteria - the factors.

Step 5) Determining the importance of SWOT sub-criteria within the third level of the model proposed in Figure 1b and formation of the matrices $w_{\text{SWOT sub-factors(local)}}$ with respect to each SWOT factor (Strengths, Weaknesses, Opportunities and Threats). Evaluation of comparative pairs of SWOT sub-criteria and determination of their local importance relative to a higher level in the model, is implemented based on Saaty's scale 1-9 [18].

Step 6) Determination of global importance of SWOT sub-criteria, i.e. the values of the weight matrix $W_3 = w_{\text{SWOT sub-factors(global)}} = w_{\text{factors}} \times w_{\text{SWOT sub-factors(local)}}$ are determined.

Step 7) The importance of each considered strategic option is determined in relation to the defined subcriteria of SWOT factors, by rating the comparative pairs of options using Saaty's scale 1-9 [18]. In this way weight matrix of importance of alternative strategies is created relative to the SWOT sub-criteria, i.e. the matrix W_4 .

Step 8) Determination of the overall importance of strategic options in the model, by forming a weighted matrix $w_{\text{alternatives}} = W_4 \times w_{\text{SWOT sub-factors(global)}}$.

3 The Prioritization of Alternative Strategies Based on the SWOT Analysis

In this paper, the ANP methodology is applied for the selection of priority of alternative strategies, based on the results of the SWOT analysis (defined SWOT factors and SWOT sub-factors), for the case of TFB for the period up to 2025.

The SWOT analysis for the TFB had been prepared for the purpose of the second round of national accreditation in 2013. The SWOT analysis was conducted through a few rounds of brainstorming, where between 70 and 80 professors and assistants had been participating [22], while the obtained results were adopted by the professional and management bodies of the TFB. The obtained results are shown in Table 2.

Step 1) Based on the results of the SWOT analysis which define the current state of the TFB, by comparing the SWOT factors: strengths, weaknesses, opportunities and threats, as well as sub-factors within each factor, the possible future development strategies of the TFB were defined until the year 2025 [7, 8, 9]. The analytic network structure of the MCDM model, which was used in this study, is shown in Figure 3.

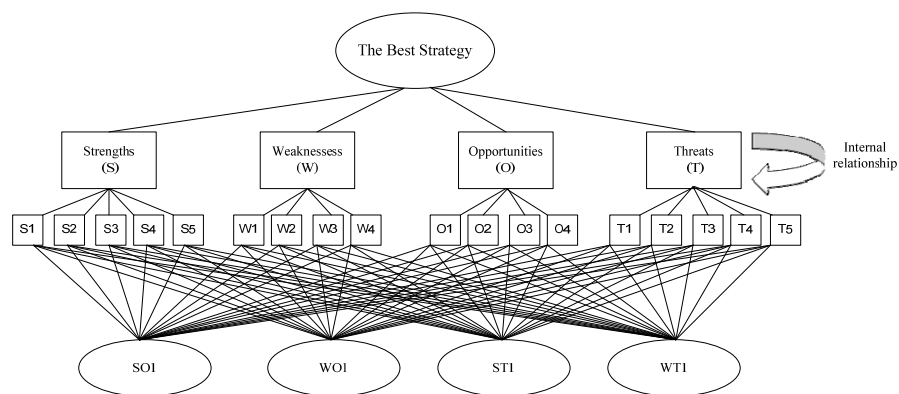


Figure 3
AHP model for the selection of the best strategy

The results shown in Table 2 identified the following strategies:

- SO_1 : Strategy for development of new markets (Providing students from the new markets in the country and abroad, as well as opening of departments outside the seat of the Faculty. TFB has previous experience with this type of activities).
- WO_1 : Strategy of the shift in the management of the Faculty (Moving from the current strategy: non-transparent management and retention of

the status quo into a transparent and aggressive strategy of making changes in order to ensure growth and development).

- ST₁: The strategy of new product development (development of new and attractive study programs which would be of interest for potential students)
- WT₁: Strategy for the development of strategic partnerships within the BU and the EU (creation of joint programs and issuance of double degrees with other units of BU and universities from the EU).

By combining the SWOT, factors and sub-factors within each factor, possible alternative SO, WO, ST and WT strategies were defined, which derive from the adopted mission statement of the TFB: "The purpose of the TFB's existence is to provide an adequate response to the needs of young generation for the higher education. The best in our field will be chosen as benchmarking partners in the realization of the educational process. Also, alternative strategies are aligned with the vision document of the TFB: "Vision of the TFB is that it becomes recognized in the educational space of South East Europe through achieving above-average results in science and education" [22].

Table 2
SWOT analysis for the TFB

Internal factors		
	Strengths (S)	Weaknesses (W)
External factors	S ₁ - Membership in UB	W ₁ - Unwillingness to change
	S ₂ - International reputation	W ₂ - Lack of additional revenue
	S ₃ - Free studies	W ₃ - Lack of leadership
	S ₄ - Good accommodation for students	W ₄ - Insufficient cooperation with the surrounding environment
	S ₅ - Online access to scientific data bases and networks	
Opportunities (O)	SO – Strategy	WO – Strategy
O ₁ - Cooperation with alumni	SO₁ – Strategy for development of new markets	WO₁ – Strategy of the shift in the management of the Faculty
O ₂ - International exchange		
O ₃ - Increased demands for quality		
O ₄ - Access to EU funds		

Threats (T)	ST – Strategy	WT – Strategy
T ₁ - Declining number of potential students	ST ₁ – The strategy of product development (development of new and attractive study programs)	WT ₁ – Strategy for formation of the strategic partnerships within the BU and EU
T ₂ - Declining living standard in Serbia		
T ₃ - Lack of students' motivation		
T ₄ - Declining level of input knowledge of new students		
T ₅ - Inconsistency of state policy		

Step 2) Based on the rankings of the expert team the importance of each SWOT factor (criteria) in the model is determined, while their internal interdependence was not considered, but only importance in relation to the objective that is set within level 1 (see Figure 1b). The resulting importance of each SWOT factor is shown in Table 3, where it can be seen that the greatest importance, based on scores of the expert team, has the SWOT factor Opportunities (42% importance).

Table 3
Pairwise comparison of SWOT groups without interdependences between them

SWOT group	S	W	O	T	Importance of the SWOT factor
Strengths (S)	1	2	1/3	1/2	0.168
Weaknesses (W)		1	1/2	1/3	0.123
Opportunities (O)			1	2	0.420
Threats (T)				1	0.289
Consistency ratio relative to the goal: CR = 0.06					

From Table 3, it follows that:

$$w_1 = \begin{bmatrix} S \\ W \\ O \\ T \end{bmatrix} = \begin{bmatrix} 0.168 \\ 0.123 \\ 0.420 \\ 0.289 \end{bmatrix}$$

Step 3) In this step, the inner interdependence of the SWOT factors is determined according to the model defined in Figure 2. Tables 4-7, show the ranks of compared pairs of SWOT factors, which are evaluated by the expert team, as well as the resulting weight vectors of internal interdependences of SWOT factors.

Table 4
Matrix of internal interdependencies of SWOT groups in relation to Strengths

Strengths (S)	W	O	T	Relative importance weight
Weaknesses (W)	1	1/5	1/7	0.075
Opportunities (O)		1	1/2	0.330
Threats (T)			1	0.595
Consistency ratio relative to the goal: CR = 0.014				

Table 5
Matrix of internal interdependencies of SWOT groups in relation to Weaknesses

Weaknesses (W)	S	O	T	Relative importance weight
Strengths (S)	1	3	6	0.635
Opportunities (O)		1	5	0.290
Threats (T)			1	0.075
Consistency ratio relative to the goal: CR = 0.09				

Table 6
Matrix of internal interdependencies of SWOT groups in relation to Opportunities

Opportunities (O)	S	W	T	Relative importance weight
Strengths (S)	1	1/3	1/4	0.124
Weaknesses (W)		1	1/2	0.517
Threats (T)			1	0.359
Consistency ratio relative to the goal: CR = 0.01				

Table 7
Matrix of internal interdependencies of SWOT groups in relation to Threats

Threats (T)	S	W	O	Relative importance weight
Strengths (S)	1	1/3	1/3	0.140
Weaknesses (W)		1	1/2	0.528
Opportunities (O)			1	0.332
Consistency ratio relative to the goal: CR = 0.05				

On the basis of the calculated relative importance weights of the SWOT factors, the inner dependence matrix W_2 is created, in the form of:

$$W_2 = \begin{bmatrix} 1 & 0.635 & 0.124 & 0.140 \\ 0.075 & 1 & 0.517 & 0.528 \\ 0.330 & 0.290 & 1 & 0.332 \\ 0.595 & 0.075 & 0.359 & 1 \end{bmatrix}$$

Step 4) Obtained relative importance weights of the SWOT factors in the inner dependence matrix W_2 , are then used for the "correction" of the initial weights of

SWOT factors which are defined by the matrix w_1 , after which importance weights of the SWOT factors become:

$$W_{\text{SWOT factors}} = W_2 \times w_1 \begin{bmatrix} 1 & 0.635 & 0.124 & 0.140 \\ 0.075 & 1 & 0.517 & 0.528 \\ 0.330 & 0.290 & 1 & 0.332 \\ 0.595 & 0.075 & 0.359 & 1 \end{bmatrix} \times \begin{bmatrix} 0.168 \\ 0.123 \\ 0.420 \\ 0.289 \end{bmatrix} = \begin{bmatrix} 0.169 \\ 0.253 \\ 0.304 \\ 0.274 \end{bmatrix}$$

Based on the newly gained priorities of interdependencies of SWOT factors, it can be noticed that there has been a significant change in the relative importance of the two SWOT factors, namely: the importance of the Weaknesses factor is now increased by 13% in the model (from 12.3% to 25.3%), while the impact of the most important SWOT factor Opportunities has now declined by 11.6% compared to the original 42% and now amounts to 30.4%.

Step 5) In this step, the local importance of SWOT sub-criteria is determined by the expert team, while the ranks of comparative pairs of the SWOT sub-criteria, defined in Table 2, are given in Tables 8-11.

Table 8
Pairwise comparisons of the SWOT sub-criterion - Strengths

Strengths (S)	S ₁	S ₂	S ₃	S ₄	S ₅	Local weights
S ₁ - Membership in UB	1	3	1/2	3	3	0.287
S ₂ - International reputation		1	1/3	3	3	0.175
S ₃ - Free studies			1	2	3	0.353
S ₄ - Good accommodation for students				1	2	0.110
S ₅ - Online access to scientific data bases and networks					1	0.075
The consistency ratio in relation to the group Strengths: CR = 0.08						

Table 9
Pairwise comparisons of the SWOT sub-criterion - Weaknesses

Weaknesses (W)	W ₁	W ₂	W ₃	W ₄	Local weights
W ₁ - Unwillingness to change	1	3	2	3	0.431
W ₂ - Lack of additional revenue		1	1/3	1/2	0.101
W ₃ - Lack of leadership			1	4	0.333
W ₄ - Insufficient cooperation with the surrounding environment				1	0.135
The consistency ratio in relation to the group Weaknesses: CR = 0.07					

Table 10
Pairwise comparisons of the SWOT sub-criterion - Opportunities

Opportunities (O)	O ₁	O ₂	O ₃	O ₄	Local weights
O ₁ - Cooperation with alumni	1	1/3	1/4	1/2	0.094
O ₂ - International exchange		1	1/2	3	0.316
O ₃ - Increased demands for quality			1	2	0.428
O ₄ - Access to EU funds				1	0.163
The degree of consistency in relation to the group Opportunities: CR = 0.04					

Table 11
Pairwise comparisons of the SWOT sub-criterion - Threats

Threats (T)	T ₁	T ₂	T ₃	T ₄	T ₅	Local weights
T ₁ - Declining number of potential students	1	2	3	3	4	0.377
T ₂ - Declining living standard in Serbia		1	3	4	3	0.291
T ₃ - Lack of students' motivation			1	4	3	0.177
T ₄ - Declining level of input knowledge of new students				1	2	0.087
T ₅ - Inconsistency of state policy					1	0.067
The degree of consistency in relation to the group Threats CR = 0.08						

Step 6) Global significance of SWOT sub-criteria is obtained by multiplying factor weights from Step 4 and Step 5 among each other, as presented in Table 12.

Table 12
Importance of criteria and sub-criteria of the SWOT analysis

SWOT groups - criteria	Importance of the SWOT group	SWOT sub-criteria	Local importance of SWOT sub-criterion	The overall importance of SWOT sub-criteria
Strengths - S	0.169	S ₁ - Membership in UB		
		S ₂ - International reputation	0.287	0.049
		S ₃ - Free studies	0.175	0.030
		S ₄ - Good accommodation for students	<u>0.353</u>	0.060
		S ₅ - Online access to scientific data bases and networks	0.110	0.019
			0.075	0.013

Weaknesses - W	0.253	W ₁ - Unwillingness to change	<u>0.431</u>	0.109
		W ₂ - Lack of additional revenue	0.101	0.026
		W ₃ - Lack of leadership	0.333	0.084
		W ₄ - Insufficient cooperation with the surrounding environment	0.135	0.034
Opportunities - O	0.304	O ₁ - Cooperation with alumni	0.094	0.029
		O ₂ - International exchange	0.316	0.096
		O ₃ - Increased demands for quality	<u>0.428</u>	0.130
		O ₄ - Access to EU funds	0.163	0.050
Threats - T	0.274	T ₁ - Reducing the number of potential students		
		T ₂ - Declining living standard in Serbia	<u>0.377</u>	0.103
		T ₃ - Lack of students' motivation	0.291	0.080
		T ₄ - Declining level of input knowledge of new students	0.177	0.048
		T ₅ - Inconsistency of state policy	0.087	0.024
			0.067	0.018

Hence it follows that:

$$W_3 = W_{\text{SWOTsub-factros(global)}} = \begin{bmatrix} 0.049 \\ 0.030 \\ 0.060 \\ 0.019 \\ 0.013 \\ 0.109 \\ 0.026 \\ 0.084 \\ 0.034 \\ 0.029 \\ 0.096 \\ 0.130 \\ 0.050 \\ 0.103 \\ 0.080 \\ 0.048 \\ 0.024 \\ 0.018 \end{bmatrix}$$

Step 7) In this step, the importance weight of each alternative strategy (SO_1 , WO_1 , ST_1 , WT_1) was determined in relation to the defined SWOT sub-criteria, which results in the matrix W_4 as following:

$$W_4 = \begin{bmatrix} 0.529 & 0.381 & 0.628 & 0.528 & 0.385 & 0.312 & 0.128 & 0.250 & 0.159 & 0.507 & 0.499 & 0.104 & 0.134 & 0.068 & 0.079 & 0.359 & 0.097 & 0.160 \\ 0.068 & 0.079 & 0.074 & 0.105 & 0.087 & 0.127 & 0.371 & 0.250 & 0.381 & 0.085 & 0.073 & 0.730 & 0.529 & 0.529 & 0.381 & 0.359 & 0.402 & 0.467 \\ 0.134 & 0.159 & 0.158 & 0.105 & 0.143 & 0.280 & 0.422 & 0.250 & 0.381 & 0.204 & 0.214 & 0.061 & 0.068 & 0.268 & 0.381 & 0.200 & 0.164 & 0.277 \\ 0.268 & 0.381 & 0.140 & 0.262 & 0.385 & 0.280 & 0.079 & 0.250 & 0.079 & 0.204 & 0.214 & 0.104 & 0.268 & 0.134 & 0.159 & 0.082 & 0.337 & 0.095 \end{bmatrix}$$

Step 8) Finally, the overall priority of the considered strategies was calculated as follows:

$$W_{\text{alternatives}} = \begin{bmatrix} SO_1 \\ WO_1 \\ ST_1 \\ WT_1 \end{bmatrix} = W_4 \times W_{\text{SWOTsub-factors(global)}} = \begin{bmatrix} 0.271 \\ 0.322 \\ 0.214 \\ 0.192 \end{bmatrix}$$

4 Discussion of Results

Application of ANP - SWOT methodology allows prioritization of the identified alternative strategies which are shown in Table 2. The priority, according to the obtained results, is defined in the following descending order: $WO_1 - SO_1 - ST_1 - WT_1$. In the SWOT criteria, according to the proposed ANP model with internal interdependence on level 2 (see Figure 1), the factors which have the most importance are the opportunities (O) - 0.304, followed by threats (T) - 0.274 and eventually the weaknesses (W) - 0.253 and strengths (S) - 0.169, while the interactions were developed between each SWOT - factor.

The sub-criteria in individual SWOT criteria with the greatest importance according to the ANP are: S (S_3 - Free studies: 0.353); W (W_1 - Unwillingness to change: 0.431); O (O_3 - Increased demands for quality: 0.428); T (T_1 - Declining number of potential students: 0.377). These facts have a dominant importance when sub-criteria S_3 and O_3 are being used to maximize the results of implemented strategies, as well as in defining a series of actions to minimize the influence of W_1 and T_1 .

Mission statement of the TFB is: "The purpose of the TFB's existence is to provide an adequate response to the needs of the younger generation for higher education. Implementation of the teaching process will be realized according to the highest standards, while the best in our field are being chosen as benchmarking partners". The mission statement of the TFB implies continued growth and development in a changing environment, which requires the use of WO_1 strategy - changing the way the Faculty is managed, from the current approach: non-transparency while maintaining the status quo, into a transparent and aggressive strategy of continuous changes in order to grow and develop using all available

resources. Due to the limited and declining market for the TFB, after the change in the management style, which is followed by the implementation of the strategy SO_1 in parallel with the strategy WO_1 or immediately after it, it is necessary to enroll students from other markets by using the adequate aggressive promotion or by opening new study centers outside the seat of the Faculty.

Above mentioned strategies WO_1 and SO_1 can provide a certain growth and development in the initial phase of further growth and development of the TFB, with a limited reach on the life cycle curve of this organization. After, reaching the limits of growth and development using the outlined WO_1 and SO_1 strategies, in order to upgrade and improve the life cycle of the TFB, the strategy ST_1 should be implemented, which, in accordance with the "mission" statement, implies defining new study programs according to the requirements of prospective students.

From the position achieved after applying strategies WT_1 - SO_1 - ST_1 , the TFB will become a desirable partner for the implementation of WT_1 - formation of the strategic partnerships with the best institutions within the BU and EU. In this way until 2025 creates a realistic chance of achieving the vision TFB: In this way, realistic chances of achieving the vision of the TFB by the year 2025 will be created: "Achieving distinctive position in the educational space of the South Eastern Europe".

In order to achieve the abovementioned goals, the following strategy implementation sequence should be applied: WO_1 - SO_1 - ST_1 - WT_1 . This will guide TFB towards the specified strategic goals which are defined in the vision. Also, it will be required to boost the importance of the sub-criteria S_3 and O_3 , with a positive impact on the realization of the strategies labeled with S and O, and to reduce, through continuous changes, negative impacts of the W_1 and T_1 in the strategies labeled with W and T.

Conclusion

The traditional SWOT analysis involves an arbitrary ranking criteria and sub-criteria independently of each other, ignoring the potential interactions between them. In order to overcome abovementioned shortcomings of the traditional SWOT analysis, an attempt was made in this paper to improve this methodology by using ANP network methodology as an upgrade of the initial SWOT matrix.

The results regarding prioritization of possible development strategies of the TFB until the year 2025, which are obtained by using the ANP-SWOT model, show that by taking into account the interactions between goals (selection of the best strategy), SWOT factors, SWOT sub factors and alternatives (possible strategies) at different hierarchical levels, as well as factors within the clusters at the same level through the ANP network model, prioritization of possible strategic alternatives can be reliably defined.

Values of the final weights in the normalized matrix of possible alternative strategies provide opportunities for prioritization of defined strategies which need

to be continuously administered during the planning period until 2025, in proportion to the progress, which is achieved with the implementation of the previous strategic alternative.

Obtained results, based on the comprehensive numerical data analysis, will serve as the starting point for benchmarking the position of the TFB in the academic scope of the region, and sustaining its future upraise and competitiveness.

References

- [1] Porter, M. F., *Comparative Strategy: Technique for Analysing Industry and Competitors*, The Free Press, New York, 1980
- [2] Gorener, A., Comparing AHP and ANP: An Application of Strategic Decision Making in a Manufacturing Company, *International Journal of Business and Social Science*, 3(11) (2012) pp. 194- 2008
- [3] Kurttila, M., Pesonen, M., Kangas, J., Kajanus, M., Utilizing the Analytical Hierarchy Process (AHP) in SWOT Analysis – a Hybrid Method and its Application to a Forest Certification Case, *Forest Policy and Economics*, 1 (2000) pp. 41-52
- [4] Hamidi, K., Delbahari, V., Formulating a Strategy for a University using SWOT Technique: A Case Study, *Australian Journal of Basic and Applied Sciences*, 5(12) (2011) pp. 264-276
- [5] Sharifi, A. S., Islamic Azad University Function Analysis with Using the SWOT Model in Order Provide Strategic Guidelines (Case study: Faculty of Humanities), *Procedia Social and Behavioral Sciences*, 58 (2012) pp. 1535-1543
- [6] Kangas, J., Kurttila M., Kajanus, M., Kangas, A., Evaluating the Management Strategies of a Forestland Estate of S-O-S Approach, *Journal of Environmental Management*, 69 (2003) pp. 349-358
- [7] Dyson, R. G., Strategic Development and SWOT Analysis at the University of Warwick, *European Journal of Operational Research*, 152 (2004) pp. 631-640
- [8] Yuksel, I., Degdeviren, M., Using the Analytical Network Process (ANP) in a SWOT Analysis- A Case Study for a Textile Firm, *Information Sciences*, 177(2007) pp. 3364-3382
- [9] Sevkli, M., Oztekin, A., Uysal, O., Torlak, G., Turkuilmaz, A., Delen, D., Development of a Fuzzy ANP-based SWOT Analysis for the Airline Industry in Turkey, *Expert Systems with Applications*, 39 (2012) pp. 14-24
- [10] Kangas, J., Pesonen, M., Kurttila, M., Kajanus, M., A'WOT: Integrating the AHP with SWOT Analysis, *Proc. 6th ISAHP 2001*, Berne, Switzerland (2001) pp. 189-198

- [11] Kajanus, M., Kamngas, J., Kurttila, M., The Use Value Focused Thinking and the A*WOT Hybrid Method in Tourism Management, *Tourism Management*, 25 (2004) pp. 499-506
- [12] Lee, J. W., Kim, S. H., Using Analytical Network Process and Goal Programming for Interdependent Information System for Project Selection, *Computers and Operations Research*, 27 (2000) pp. 367-382
- [13] Gorener, A., Toker, K., Ulucay, K., Application of Combined SWOT and AHP: A Case Study for a Manufacturing Firm, *Procedia – Social and Behavioral Sciences*, 58 (2012) pp. 1525-1534
- [14] Wicramasinghe, V., Takano, S., Application of Combined SWOT and Analytic Hierarchy Process (AHP) for Tourism Revival Strategic Marketing Planning: A Case of Sri Lanka Tourism, *Journal of the Eastern Asia Society for Transportation Studies*, 8 (2009) pp. 1-16
- [15] Khamseh, A. A., Fazauei, M., A Fuzzy Analytical Network Process for SWOT Analysis (Case Study: Drug Distribution Company) *Technical Journal of Engineering and Applied Sciences*, 3 (18) (2013) pp. 2317-2326
- [16] Kheirkhah, A., Babaeianpour, M., Bassiri, P., Development of a Hybrid Method Based on Fuzzy PROMETHEE and ANP in the Framework of SWOT Analysis for Strategic Decisions, *International Research Journal of Applied and Basic Sciences*, 8 (4) (2014) pp. 504-515
- [17] Chang, H. H., Huang, W. C. (2006) Application of a Quantification SWOT Analytical Method. *Mathematical and Computer Modelling*, 43: 158-169
- [18] Saaty, T. L., *The Analytical Hierarchy Process*, New York: McGrawHill, 1980
- [19] Saaty, T. L., Vagras, L. G., *Models, Methods, Concepts and Applications of the Analytical Hierarchy Process*, Boston, MA: Kluwer Academic Publishers, 2001
- [20] Saaty, T. L., Takizawa, M., Dependence and Independence from Linear Hierarchies to Nonlinear Networks, *European Journal of Operational Research*, 26 (1986) pp. 229-237
- [21] Stanujkić, D., Djordjević, B., Djordjević, M., Comparative Analysis of some Prominent MCDM Methods: A Case of Ranking Serbian Banks, *Serbian Journal of Management*, 8 (2) (2013) pp. 213-241
- [22] Savić, M., Djordjević, P., Nikolić, Dj., Mihajlović, I., Živković, Ž., Bayesian Inference for Risk Assessment of the Position of Study Program within the Integrated University: a Case Study of Engineering Management at the Technical Faculty in Bor, *Serbian Journal of Management*, 9(2) (2014) pp. 231-240

Accelerated Half-Space Triangle Rasterization

Péter Mileff, Károly Nehéz, Judit Dudra

University of Miskolc, Department of Information Technology, Miskolc-Egyetemváros, 3515 Miskolc, Hungary, mileff@iit.uni-miskolc.hu

University of Miskolc, Department of Information Engineering, Miskolc-Egyetemváros, 3515 Miskolc, Hungary, aitnehez@uni-miskolc.hu

Bay Zoltán Nonprofit Ltd. for Applied Research, Engineering Division (BAY-ENG), Department of Structural Integrity and Production Technologies, Iglói út 2, 3519 Miskolc, Hungary, judit.dudra@bayzoltan.hu

Abstract: Since the last decade, graphics processing units (GPU) have dominated the world of interactive computer graphics and visualization technology. Over the years, sophisticated tools and programming interfaces (like OpenGL, DirectX API) greatly facilitated the work of developers because these frameworks provide all fundamental visualization algorithms. The research target is to develop a traditional CPU-based pure software rasterizer. Although currently it has a lot of drawbacks compared with GPUs, the modern multi-core systems and the diversity of the platforms may require such implementations. In this paper, we are dealing with triangle rasterization, which is the cornerstone of rendering process. New model optimization as well as the improvement and combination of existing techniques have been presented, which dramatically improve performance of the well-known half-space rasterization algorithm. The presented techniques become applicable in areas such as graphics editors and even computer games.

Keywords: half-space rasterization; bisector algorithm; software rendering

1 Introduction

The history of computer visualization goes back several decades, but its importance has grown more significantly in recent years. More advanced graphical features are dominant almost everywhere these days, from traditional desktop computers to mobile and embedded devices. This is a result of a long development process, which was mainly induced by the appearance of graphical processors. While in early computer visualization only slow CPUs were used to perform every stage of the entire rasterization process, nowadays graphics cards (target hardware) have taken over this role. The main development directions of the area are inspired and usually controlled by the professional computer game and media industry and the increasing demands for CAD/CAM systems.

However, we should not forget about the recent development of CPUs. Modern CPUs have many advanced features due to multi-core technology and the extended instruction set (e.g. SSE, AVX). Besides the development of the instruction set and the increase of central cores, another important result is the appearance of DDR4 type memories in 2014. Although these types are one order of magnitude faster than older DDR3 RAMs, still cannot compete with the DDR5 type memory equipped in modern video cards. Nevertheless it is a significant step forward to improve the speed of memory operations. The continuous development of CPU technology encourages software developers to reconsider the structure and logical model of their existing graphical applications. The usage of a multi-thread game engine model is essential for today's AAA-type computer games. Developers of the most advanced graphics engines (Unreal Engine, CryEngine, Frostbite, etc.) have already recognized the potential of these new opportunities. The question may arise, whether it makes sense to deal with CPU-based solutions if powerful GPUs are available today. The answer has already been given by leading video game developers. Some modern games apply the CPU to perform specific tasks to reduce GPU load. Software occlusion culling - where the CPU is used to render polygons to occlusion buffer rather than the GPU - is a good example of this hybrid approach. Several well-known games (e.g. KillZone 3, Battlefield 3, etc.) and game engines (e.g. CryEngine) apply similar technologies because CPUs have no latency problem of occlusion queries. Rasterization is usually done in small resolution (e.g. Battlefield uses 256×114) and occlusion testing can be done using a hierarchical software z-buffer [12] [19].

Using all this as a starting point, the central question that motivates this paper is how to improve one of the fundamental visualization algorithms, i.e. a polygon fill based rasterization on CPUs. Our objective is to propose algorithms and extensions for the half-space rasterization model, which can serve as the basis of a graphics engine applying more complex, possibly a hybrid graphics pipeline using CPU for specific tasks.

2 Related Works

Software based rendering prospered between the end of the 90s and the beginning of the 2000's. Due to the appearance of GPUs, only a few papers (although an increasing number of) that involve the CPU in the visualisation process have been published in recent years. Most of these papers discuss general display algorithms or shader oriented GPU specific solutions which cannot be applied on CPUs in their original form.

During the early years of the rendering (1996-1998) ID Software and Epic Games achieved remarkable results in the area of modern software based computer graphics. Both companies have become famous for their high performance and complex graphics engine offering high quality visualization solutions (e.g. colored

lighting, shadowing, volumetric lighting, fog, pixel accurate culling, etc.). These engines were optimized for the Intel Pentium MMX processor family and their rendering system was based on the scanline-based polygon filling approach applying several additional technologies (e.g. BSP tree), so as to provide high performance. After the continuous spreading of GPU-based rendering, software rendering was pushed more and more into the background. Nevertheless, some great results, e.g. Pixomatic Renderer developed by Rad Game Tools and the Swiftshader [2] by TransGaming were achieved. Both products are fully DirectX 9 compatible, very complex and highly optimized taking advantage of multi-core threading possibilities of modern CPUs. Their pixel pipeline can continuously modify itself adapting to the actual rendering tasks. Since these products are all proprietary, the details of their architectures are not available for the general public. By developing DirectX, Microsoft provided the basis for the spread of GPU technologies, and it also developed a software rasterizer called WARP [3]. The renderer is capable of taking advantage of multi-threads and in some cases it is even able to outperform low-end integrated graphics cards.

In 2008, based on problem and demand investigations, Intel aimed to develop its x86 based video card within the Larrabee project [5]. In a technological sense, the card was a hybrid of the multi-core CPUs and GPUs. Its purpose was to provide x86 cores-based, fully programmable pipeline with 16-byte-wide SIMD vector units. The new architecture made it possible for graphic calculations to be programmed in a more flexible way than GPUs with an x86 instruction set.

Other researchers proposed optimization of the triangle traversal algorithms, and new rasterization models have been introduced. As a part of current results, Hengyong Jiang et al. [13] proposed a midpoint triangle rasterization traversal algorithm, which reduces the number of traversal points and improves the efficiency of graphics acceleration. Their approach is demonstrated by an FPGA. Pablo et al. investigated and compared three different triangle traversal algorithms (Box, Zig-zag, Hilbert Curve based) in terms of performance and they simulated them in Matlab using ModelSim [22]. Their experimental results show that important area-performance trade-offs can be met, when implementing key image processing algorithms in hardware. Chih-Hao Sun et al. [15] demonstrated an edge equation based tile-scan triangle traversal algorithm. In their solution, the basic functions of parameter interpolation and rasterization can be executed with a universal shared hardware to reduce the cost of rendering. By hardware sharing and architecture design techniques of pipelining and scheduling, their algorithm can meet real-time requirements for graphics applications at reasonable hardware costs. An entirely new approach was presented by Olano et al., which is a simplified solution for triangle scan conversion applying 2D homogeneous coordinates for fast real-time rendering [17]. Their solution avoids costly clipping tests and can render true homogeneous triangles significantly faster than previous implementations. As a part of the new generation of parallel algorithms, Zack Bethel outlined a modern, multi-thread tile based software rendering technique

using a block-based half-space theory where only the CPU is used for calculations, which led to performance improvements [1]. The FreePipe Software Rasterizer [9] focuses on multi-fragment effects, where each thread processes one input triangle, determines its pixel coverage, and performs shading and blending sequentially for each pixel. Due to the evolution of GPGPU, the idea of performing software rendering by GPGPU has been raised several times. NVidia proposed and investigated an efficient CUDA-based rendering model [6], the performance of which is a factor of 2-8x compared to the hardware graphics pipeline.

In recent years, major companies in the game industry have recognized the potential of the CPU again [11]. Their game engine can delegate certain visualization tasks to the CPU. In the game Battlefield 3 a SPU-based deferred shading model was developed [12], where the objective was to use SPUs to distribute shading work and offload GPU. In 2011, the game Killzone 3 supports complex occluded environments. To cull non-visible geometry early in the frame, the game uses PlayStation 3 SPUs to rasterize a conservative depth buffer and performs fast synchronous occlusion queries against it [16]. The topic of software occlusion culling has been investigated also by Intel Software in the paper [18]. It presents a Killzone-like solution, but it is built upon an x86 basis and optimized for SSE Streaming extensions.

Thus, as recent findings show, CPU-based approaches put an emphasis on again to increase flexibility and performance. This paper investigates the optimization and extensions of the triangle traversal and filling algorithm.

3 Basics of Rasterization

The aim of computer visualization is to display pixel sets (e.g. 2D image or projected 3D objects) on the screen. The type of rendering algorithm or the procedure of a presentable element largely depends on the applied hardware or software based visualisation model. Rasterization is a very intensive process computationally, especially when the visual element also contains an alpha channel [7] [21].

Several approaches have been developed to represent shapes in memory, but nowadays the most prevalent and most widely applied object representation is the polygon mesh. During the modelling process the object is usually divided into convex polygons, like triangles. Still, the rendering performance largely depends on the applied rasterization algorithm. Although several different solutions have been developed (Ray tracing, Volume Rendering, etc.), currently GPU manufacturers use triangle filling based models in real-time visualisation. This method allows significantly faster rendering than for example ray based algorithms.

3.1 Triangle-based Filling

In a classical sense filling is performed pixel by pixel, so the inner iteration and the various calculations need to be executed many times. Although the vertex mapping and the traversal of the process seem to be simple, the filling performance largely depends on the implemented algorithm and its optimization level. We can state that it is a huge challenge to implement an algorithm, which takes the possibilities of the modern CPU hardware into account and being highly optimized at the same time. The parts of the filling model affect each other in a complex form. The smallest change in the iteration logic can result in up to more than a 10% performance difference.

Nowadays there are two widespread triangle filling algorithms: the scanline and the half-space based algorithms. The main idea of the previous scanline approach was to walk triangle line by line (scanline) from top to bottom. Each row represents a line, whose starting and ending points are the intersections of the triangle sides and the scanline along the x axis. The end points of the line can be calculated incrementally using the slope values of the edges.

The scanline algorithm is widely used and can be optimized (e.g. s-buffer), but it is difficult to adapt it to current hardware opportunities. The rasterization is performed line by line, which has several unpleasant consequences for both hardware and software implementation. One of these problems is that the algorithm is asymmetric in x and y directions. In the case of thin triangles, the performance may significantly vary between the horizontal and vertical orientation. The outer scanline loop is serial, which is not favourable for hardware implementation. The inner loop, which is responsible for scanline iteration is not so SIMD-friendly because of the different line lengths. This makes the algorithm orientation-dependent. Processing several lines at the same time for some reason (e.g. Mipmapping, Multisampling) would mean further complications for calculations. To sum up, the solution is hard to apply for parallel processing of lines and pixels. A better solution would be if pixels were processed in 2x2 blocks (quads), the expected increase in the performance would be significant.

Next, the paper focuses on the other approach, the half-space based triangle traversal, where the basic algorithm and further optimization points are presented which improves performance to a large extent.

4 Half-Space Rasterization

The name of this model is not unified in literature; some sources refer to it as a point in a triangle, a bounding box or a half-plane algorithm. The basic idea of the model originates from the polygon convexity: the interior of a convex polygon formed by n number of edges can be defined as the intersection of the n half spaces.

For triangles, the three half-planes clearly define the inner area.

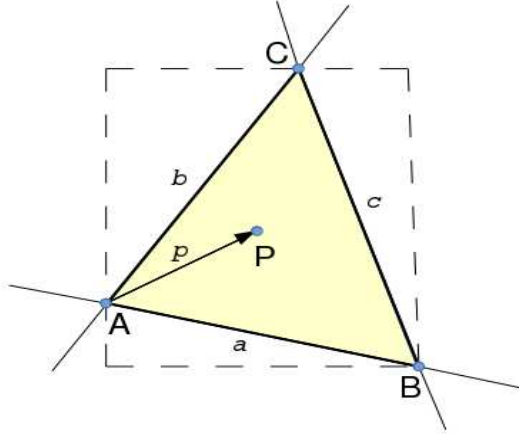


Figure 1

Triangle, defined by half-planes

Several mathematical approaches describe the inner pixels of the triangle and they all examine the question of how to describe edges. In the following a ‘Perp Dot’ product-based model is presented. The Perp Dot product [20] is the two-dimensional equivalent of the three dimensional cross-products. It is a product between two vectors in two dimensions and is obtained by taking the dot product of one vector with the perpendicular one of the other. The perpendicular vector is simply a vector at right angles to the vector it is based on with the same magnitude. The definition of the formula for two-dimensional vectors and applying it to a, b vectors is the following:

$$\text{perpDotP}(a, b) = a^\perp \cdot b = a_x b_y - a_y b_x = \begin{vmatrix} a_x & a_y \\ b_x & b_y \end{vmatrix}. \quad (1)$$

The result is a scalar, which has specific properties:

- $\text{perpDotP}(a, b) = 0 \rightarrow a, b$ are parallel
- $\text{perpDotP}(a, b) > 0 \rightarrow b$ is counterclockwise from a
- $\text{perpDotP}(a, b) < 0 \rightarrow a$ is counterclockwise from b

These properties are useful to determine whether a pixel is inside the triangle or not. In the case of a P point, the product needs to be calculated with all the three edges and to check its sign,

$$c_1 = \text{perpDotP}(a, p), c_2 = \text{perpDotP}(b, p), c_3 = \text{perpDotP}(c, p).$$

The final point-triangle containment relation depends on the prior knowledge of the triangle:

- if the vertices A, B, C are given in clockwise order, then P is inside the triangle if $c_1 > 0 \ \&\& \ c_2 > 0 \ \&\& \ c_3 > 0$
- if the vertices A, B, C are given in counter-clockwise order, then P is inside the triangle if $c_1 < 0 \ \&\& \ c_2 < 0 \ \&\& \ c_3 < 0$
- if the order of the vertices A, B, C is unknown, then P is inside the triangle if

$$(c_1 > 0 \ \&\& \ c_2 > 0 \ \&\& \ c_3 > 0 \ || \ c_1 < 0 \ \&\& \ c_2 < 0 \ \&\& \ c_3 < 0)$$

4.1 The Simple Filling Approach

Based on the equations above, a simple filling algorithm can be formulated. First of all a bounded area is needed to specify which set of pixels is required to travel. We can use the area of the render target (e.g. screen buffer), but it requires a lot of unnecessary iterations. A better solution is if we use the axis aligned bounding box of the triangle. The pseudo code of the algorithm:

```

Calculate triangle bounding box (minX,minY,maxX,maxY);
Clip box against render target bounds(minX,minY,maxX,maxY);
Loop i=minX to maxX
  Loop j=minY to maxY
    P = P (i,j);
    c1 = perpDotProduct( $\overline{AC}$ ,  $\overline{AP}$ );
    c2 = perpDotProduct( $\overline{BC}$ ,  $\overline{BP}$ );
    c3 = perpDotProduct( $\overline{CB}$ ,  $\overline{CP}$ );
    if ( c1 >= 0 and c2 >= 0 and c3 >= 0 )
      renderPixel(P);
  end
end

```

As the first step, the algorithm calculates the axis-aligned bounding box of the triangle and performs the clipping according to the bounds of the render target. In the second part of the process, all the points of the bounding box are crawled. Applying the above Perp Dot product formula the pixel-triangle containment relation can be determined.

Firstly, it is important to emphasize that this model is simple and easy to understand. Due to the bounding box the lines have the same length. Therefore, this approach is more SIMD friendly than the classical scanline algorithm. Taking advantage of the symmetry of the lines the model is much more parallel-friendly, the block or tile-based approach can be used more effectively. The processing logic is highly customizable for the CPUs used currently.

Although the basic algorithm seems simple, the model is not recommended in practice for several reasons (slow performance, lack of sub-pixel precision and filling rules). The main disadvantage is that the solution also travels unnecessary pixels, which means a lot of superfluous iterations in case of large and elongated triangles. The performance at these types of triangles will be significantly lower.

In the following, new approaches and optimizations will be presented, which dramatically improve the performance of the basic algorithm making it suitable for real-time games and other graphics-intensive applications.

5 Optimization Considerations

Developing a fast rasterizer is a complex task, high-level and multi-layered (code and logic level) optimization is required. In order to achieve really good results, full graphics pipeline optimization should be performed. However, the present paper focuses on improving the rasterization performance. The rasterization process constitutes the dominant part of the total performance requirement of the image synthesis [21]. Thus, any kind of performance improvement (even eliminating a division) affects the final result significantly.

Optimization should be performed on the basis of Michael Abrash's idea: 'Assume nothing' [4]; any logical considerations and their results can only be justified by measurements; the assumptions themselves are not acceptable. During the optimization process two main goals can be set for the filling algorithm:

- The key issue is the acceleration of the calculations and iterations because the basic algorithm is not efficient in CPU time.
- A more effective solution is required to traverse the pixels of the bounding box

5.1 Incremental Approach

One of the main problems of the simple filling model is that the pixel by pixel filling performed in the double iteration loop is extremely expensive computationally. To optimize these loops, let us start from the calculation of the Perp Dot product determinant. Based on the formula (1), the relationship of a point P and edge AB can be described as follows:

$$\begin{vmatrix} B_x - A_x & P_x - A_x \\ B_y - A_y & P_y - A_y \end{vmatrix} = (B_x - A_x)(P_y - A_y) - (B_y - A_y)(P_x - A_x). \quad (2)$$

Performing the multiplications and by rearranging the factors (2) can be written as:

$$F_{01}(P) := (A_y - B_y)P_x + (B_x - A_x)P_y + (A_x B_y - A_y B_x). \quad (3)$$

The resulting equation is called edge function. Examining the factors, if the vertex positions are constants, the function is linear for P . The other two edge functions applying the same transformation are the following:

$$F_{12}(P) := (B_y - C_y)P_x + (C_x - B_x)P_y + (B_x C_y - B_y C_x), \quad (4)$$

$$F_{20}(P) := (C_y - A_y)P_x + (A_x - C_x)P_y + (C_xA_y - C_yA_x). \quad (5)$$

If a triangle is counter-clockwise, then every edge function is positive for P , and the point is inside the triangle. F_{01} equation can be simplified by introducing additional constant:

$$I_{01} := A_y - B_y, \quad (6)$$

$$J_{01} := B_x - A_x, \quad (7)$$

$$K_{01} := A_xB_y - A_yB_x. \quad (8)$$

By substituting:

$$F_{01}(P_x, P_y) = I_{01}P_x + J_{01}P_y + K_{01}. \quad (9)$$

During the iteration we travel row by row: move one pixel to the right, then one pixel up or down. Because F_{01} is linear, axis-aligned unit steps can be calculated in both directions by using the introduced constants (6) (7) (8):

$$F_{01}(P_x + 1, P_y) - F_{01}(P_x, P_y) = I_{01}, \quad (10)$$

$$F_{01}(P_x, P_y + 1) - F_{01}(P_x, P_y) = J_{01}. \quad (11)$$

If we move one pixel to the right, I_{01} should be added to the edge function. If we move one pixel up or down then J_{01} is the additive tag. Following the above logic, the pseudo implementation of the iteration-based and modified filling algorithm is as follows:

```

Calculate triangle bounding box (minX, minY, maxX, maxY);
Clip box against render target bounds (minX, minY, maxX,
maxY);
Const.: I01, J01, K01, I02, J02, K02, I03, J03, K03, F01, F02, F03

Cy1=F01; Cy2=F02; Cy3=F03;
Loop j=minY to maxY step 1
  Cx1=Cy1; Cx2=Cy2; Cx3=Cy3;
  Loop i=minX to maxX step 1
    if (Cx1 > 0 and Cx2 > 0 and Cx3 > 0 )
      RenderPixel(P);
      // Inner loop increments
      Cx1 -= I01; Cx2 -= I02; Cx3 -= I03;
    end
    // Outer loop increments
    Cy1 += J01; Cy2 += J02; Cy3 += J03;
  end
end

```

5.2 A Block-oriented Half-Space Rasterization Model

The performance requirement of rasterization is caused by travelling all the pixels of the bounding box because the box also contains pixels, which are outside the triangle. Therefore, it is expedient to propose an extension of the basic model, by

which these pixels can be skipped. However, it is important to note that only those logical approaches can be applied which can be quickly evaluated and executed. Any operation executed in a double iteration loop and storing additional intermediate states has a negative impact on the performance. In such cases the cost of travelling unnecessary pixels will be less than the execution of control logic. In the following an efficient block-based traversal algorithm will be presented.

The unnecessary calculations can really be reduced, if the traversal process is performed on larger pixel sets. To achieve this, the triangle bounding box should be divided into squares which will be the basic units of the traversal logic. Figure 2 shows the essence of the block-based traversal method:

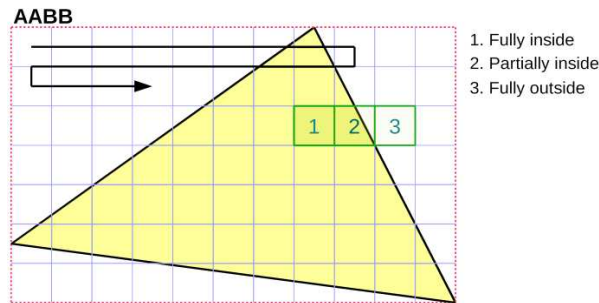


Figure 2

Block-based triangle covering

The smallest unit of iteration is a block. The algorithm is performed from top to bottom, row by row and the top-left block can be chosen as the starting point of the traversal. This approach provides an opportunity for the pixel-triangle containment relation to be performed at a block level. It is sufficient to investigate only the four corners of the square. On this basis, there are three cases:

1. every corner of the block is outside the triangle;
2. the block overlaps the triangle;
3. the block is completely inside the triangle.

The first case is the most favourable because in this case there is no need for additional calculation on writing pixels into the framebuffer. The block can be skipped entirely. The second case is the worst, where the pixel-triangle containment calculation needs to be done for each pixel in the block, and according to the results the color of the pixel should be calculated. This part of the algorithm is essentially the same as the basic, pixel-level filling. Case 3 is also favourable for traversal. Since each pixel of the block falls within the triangle, there is no need to calculate the edge functions per pixel or do the containment verification. Only the color of the pixels should be determined. The iteration also allows an additional supplement. When the edge of a triangle is reached in a row,

we can move straight to the next row because only empty blocks can be found from this position.

This traversal logic is simple, the block-level iteration does not require complex calculations: only a state should be stored and an additional condition check is required. Practical experience shows that the extra computing capacity for executing the logic is slightly smaller than the cost of checking blocks to the end of the row. However, the same solution cannot be applied for improving performance of the pixel level processing inside a block type 2. Although it seems appropriate, it degrades the performance.

The following pseudo code shows the logic of the algorithm:

```

Calculate triangle bounding box (minX, minY, maxX, maxY);
Clip box against render target bounds (minX, minY, maxX, maxY);
Loop j=minY to maxY step=q
  Loop i=minX to maxX step=q
    // Block corners
    C1x=i; C1y=(i+q-1) C2y=j; C2x=(j+q-1);
    if C1x, C2x, C1y, C2y all outside the triangle then
      continue;
    // Fully covered blocks
    if C1x, C2x, C1y, C2y all inside the triangle then
      RenderBlock(i,j,q);
      continue;
    end
    // Partially covered blocks
    Loop k=j to q step=1
      Loop l=i to q step=1
        if pixel(k,l) inside the triangle
          RenderPixel(k,l);
        end
      end
    end
  end
end
end

```

The key element of the algorithm is the iteration block size. If the block size is too small, then the block-level calculations require more computing capacity than checking only simple pixels. If large block size is chosen, then there will be fewer triangles which fulfil cases 1 or 3. Therefore, the block-based traversal does not result in any performance improvement. Practical measurement experience shows that an 8x8 block size typically provides accurate results.

The efficiency of the process can significantly be affected by the orientation of the triangles. If the presented triangle mesh consists of large polygons, or the camera view is close to a polygon, then the vast majority of the blocks will be in the triangle (case 3). In this case, rasterization can be very fast, especially when applying a SIMD instruction set in performance critical parts. However, if narrow triangles are dominating, there will be no notable speed improvement because the number of overlapping blocks is increasing (case 2).

Investigating modern computer games, it can clearly be seen that mainly pixel-level effects are dominating (bump/normal/parallax mapping, etc.). Wherever

possible, programmers apply screen space pixel transformations instead of increasing the number of polygons. Therefore, the block-based approach is well applicable.

5.2.1 Benefits of the Block-based Model

The advantage of this approach is that during the process a lot of pixels outside the triangle can be excluded from rasterization. The solution reduces the number of iterations in the bounding box, which can fundamentally be a performance-enhancing factor. Certain operations like perspective correct texture mapping or mipmapping can be performed in larger units applying linear interpolation. It is also favourable for visibility determination algorithms. In addition, a hidden advantage of the block-based approach is that it allows a more localized memory access and efficient CPU cache usage because data are located close each other. The reduction of cache misses represents a further performance improvement [10]. Finally, it should not be forgotten that the block-based model is much more parallel-friendly. Due to the distribution of operations to several threads, the block as a larger unit has important advantages.

5.3 An Adaptive Half-Space Rasterization Model

The presented pixel- and block-level models can be applied in any triangle orientation. However, their efficiency is unbalanced. The algorithm handling only pixels is less efficient in the case of large triangles and the performance of the block-based approach is not sufficient in the case of thin triangles. Another performance-enhancing factor could be, if an adaptive model was used in accordance with the characteristics of the triangle and changes were made in the applied rasterization method dynamically. The precondition for this is to determine the characteristic of the triangle.

It is expedient to define two groups: triangles, which are narrow and those which correspond to the block-oriented approach. The orientation can be defined by introducing a metric based on the bounding box of the triangle:

$$orient_{\Delta} = \frac{BB_{maxx} - BB_{minx}}{BB_{maxy} - BB_{miny}} \quad (12)$$

The definition of the formula is intentionally simple because rasterization setup costs should be kept at a minimum level. The orientation is a positive number. Its value is 1 if the bounding box is a rectangle. If the width is dominant, the value of the ratio increases and if the height is dominant, the value will be reduced. We can state that triangles with a ratio close to one are ideal for the block-based model, if their size reaches at least one block along the x or y axis. The more we move away from this value the more advisable it is to use the pixel-based algorithm. On this basis, a range for switching between the two methods can be determined by experience. Experience shows that it is triangles with orient values of 0.4 -1.6 that

should be rasterized with the block-based approach, in all other cases the pixel-based algorithm is recommended.

The adaptive algorithm is especially advantageous since it is always the preferred model that is applied of the two. Therefore, the rasterization performance can be significantly increased.

5.4 Problem of Empty Blocks

Although the block-based approach significantly helps to travel the empty area of the bounding box faster, the number of traversed empty blocks is still significant and superfluous. The following example illustrates the problem well:

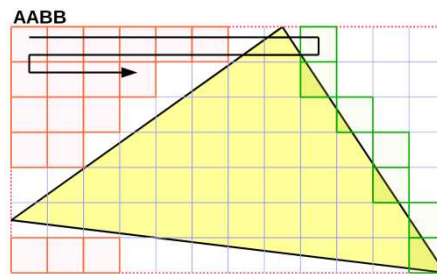


Figure 3

Block-based covering with empty block coloring

In Figure 3, the orange color indicates the traversed empty blocks. Since the algorithm always starts the investigation of the rows from the left of the bounding box, it travels several empty blocks before reaching the edges. Although these calculations are computationally less intensive than reading or writing pixels, but in the case of a complex scene a significant amount of unnecessary calculations are involved. In the above example (Figure 3), there are 18 completely empty blocks of all the 84, which is 21% of all the blocks. To reduce the traversal of the empty area, several approaches have been developed, where the ZigZag [22] and Backtrack [7] are the most widely known algorithms. Both algorithms perform rasterization at pixel level and their main characteristic is that they exclude unnecessary traversal at only one side of the triangle. However, on the other side, when stepping to the next row, they are not able to skip every empty pixel in many cases. In addition, the traversal direction changes row by row, which requires storing more states and increasing the number of conditions (CPU branching can be slow if statements are unpredictable).

5.5 A Block-based Bisector Half-Space Rasterization Model

In the following a block-based algorithm is presented which aims to minimize the number of traversed empty blocks and thus to reduce unnecessary calculations. Figure 4 explains the essence of the traversal.

The basic idea is that the problem of empty-block minimization can effectively be solved if the traversal starts from a common point inside the triangle and moves in two directions. The filling is divided into two directions along the y axis of the bounding box and performed from top to bottom, bottom to top and to half of the box.

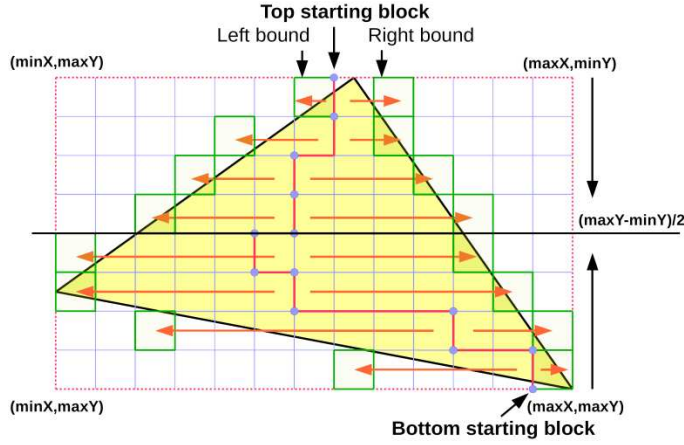


Figure 4

Triangle traversal logic of the Bisector algorithm

The bidirectional traversal at row level starts from the inside out until an empty block is reached. The starting block of the next row is calculated from the first two blocks found outside the edges at the previous row:

$$\text{Start block} = \text{Left bound} + (\text{Right bound} - \text{Left bound}) / 2. \quad (14)$$

It is important to note that, if the starting block does not fall on a block boundary, then rounding is also required. The division by two can significantly increase the load. Therefore, in implementations it is expedient to use fixed-point arithmetic because starting block can be calculated without division and with block boundary alignment:

$$\text{Start block} += (\text{Left} + (\text{Right} - \text{Left}) \gg 1) \& \sim(q - 1). \quad (15)$$

The traversal to the right direction starts from this block and the left starts filling from the previous block. The first row is special because there is no information available from the previous row. In this case, the x coordinate of the topmost vertex aligned to block boundary can be used as a starting block because it is inside the triangle. The initial state cannot be affected if two vertices of the

triangle are located at the top of the bounding box. Selecting either of them provides accurate results because these points are located at the edges of the bounding box.

Because of the block nature of the algorithm, due to the nature of the triangle edges, the iteration to the left will not find any fillable block if only one block can be filled in the first row. In this case, the value of the 'Left bound' takes the boundary coordinates of the bounding box and based on (15) the next row will definitely start at an outside block until the traversal reaches the interior of the triangle along the y axis. This is a serious problem because it increases the traversed empty blocks. To remedy this, the algorithm requires a further extension: the value of the 'Left bound' should only be changed if any inside block is found during the left directed traversal. Otherwise, its position should not be updated, so the next row also starts from the block of the previous 'Left bound', from the x coordinate of the topmost vertex.

The next key element of the algorithm is the requirement of the bottom-up traversal. Figure 5 illustrates the basic problem performing only the top-bottom crawling.

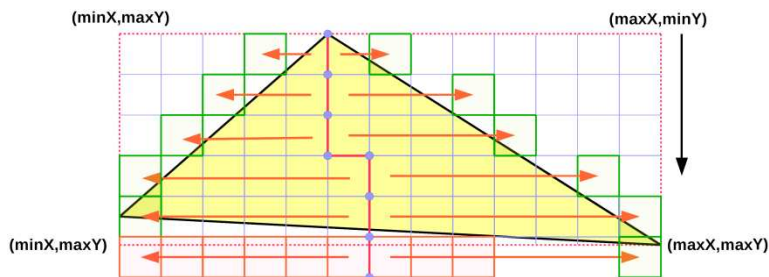


Figure 5

The inefficiency problem of the top-bottom only traversal

In cases when an edge of the triangle is very flat, the calculation of a new middle block at the next row will result in an outside block (last row in Figure 5). The best case for the traversal would be when the middle block arrives at the corner of the bottom edges, but it cannot be guaranteed because the middle point is defined by the previous row. Therefore, the number of traversed empty blocks can increase.

The solution is offered by the bottom-up traversal of the lower part of the triangle. As the starting block can exactly be determined from the x coordinate of the bottommost vertex at this time, the traversal of the empty blocks seen in Figure 5 is completely eliminable like in Figure 4.

Although the solution appears complex, it does not require any complicated calculations and the control logic is simple. The algorithm effectively minimizes

the traversal of the empty blocks. The pseudo-code of the full algorithm is as follows:

```

Calculate triangle bounding box (minX, minY, maxX, maxY);
Clip box against render target bounds (minX, minY, maxX, maxY);

Find Topmost vertex, calculate leftPoint and rightPoint
halfY = minY + ((maxY - minY) >> 1) & ~(q - 1);
BlockSize = q;
Loop j=minY to (halfY + BlockSize) step=BlockSize
    midPoint = leftPoint + ((rightPoint-leftPoint) >> 1) & ~(BlockSize - 1);
    x = midpoint;
    Loop k=0 to 2 step=1
        Loop x to (q > 0 ? x < maxX : x > minX - BlockSize) step=q
            // Block corners
            C1x=x; C1y=(x+BlockSize-1) C2y=j; C2x=(j+BlockSize-1);
            if C1x, C2x, C1y, C2y all outside the triangle then
                continue;
            // Fully covered blocks
            if C1x, C2x, C1y, C2y all inside the triangle then
                RenderBlock(x,j,BlockSize);
                continue;
            end
            RenderPartiallyCoveredBlock(j,x, BlockSize);
        end
        q = -q;
        if k==0 then rightPoint = x - BlockSize;
        else leftPoint = x + BlockSize;
        x = midpoint - BlockSize;
    end
end

Find Bottommost vertex, calculate leftPoint and rightPoint
Loop j=maxY & ~(BlockSize-1) to halfY step=-BlockSize
    Repeat the above part
end

```

6 Practical Experience and Results

To evaluate the above rasterization models from a practical point of view, a single-threaded pipeline architecture was developed. Although it is possible to apply distributed approaches in several parts of the pipeline, this paper only focuses on accelerating single-threaded pixel operations.

6.1 The Test Environment

The sample programs were written in C++ applying the GCC 4.8.1 compiler and the measurements were performed by an Intel Core i7-870 2.93 GHz CPU in a 64 bit Windows 7 environment. The implemented pipeline does not contain any hand-optimized SIMD code parts, only the compiler-optimized code was applied. The chosen screen resolution and color depth were 1024x768x32 in windowed mode and the used hardware for the test was an ATI Radeon HD 5670 with 1 GB of RAM. The software framebuffer was defined as a 32 bit unsigned integer array aligned to 16 bytes. This made an effective pixel handling possible, which stored the four components of a pixel together [21] to make memory operations faster. The prototype application used a software z-buffer and backface culling to solve visible surface determination, but texture mapping had not been implemented yet. To illuminate objects, the Lambertian reflectance was applied and rasterization used a top-left convention for filling.

6.2 Benchmark Results

During benchmarking, several different test cases were prepared. Each of them represented a special group of tasks frequently occurring in practice. The measured results are cumulated values calculated from the average frame rates (FPS) during a 20 sec running period. The distance of the models from the camera affects the performance largely. When an object is farther away, usually many small triangles should be drawn. However, getting closer to the camera, the projection of the polygons will be larger and the number of fillable pixels is increasing. Test Cases for benchmark:

Case 1: low poly model located farther from the camera.

Case 2: low poly model located close to the camera. The polygons cover about the 80% of screen pixels.

Case 3: high poly model built from small triangles (head) located farther from the camera

Case 4: high poly model built from small triangles (head) located close to the camera. The polygons cover the entire screen.

Case 5: medium poly model (statue) cover about the 80% of screen pixels. The model contains both small and large triangles.

Table 1 summarizes the results achieved by different types of rasterization algorithms. The test results show that the basic half-space algorithm is proved to be the slowest in every case as expected. Through its simplicity, it does not contain any optimizations. The incremental based approach can be regarded as a transition, its performance converges to the other, better solutions. If we examine the block-based model, it can be seen that during the tests C3 and C4 its performance was not satisfactory. The main reason for this is the nature of the scenes.

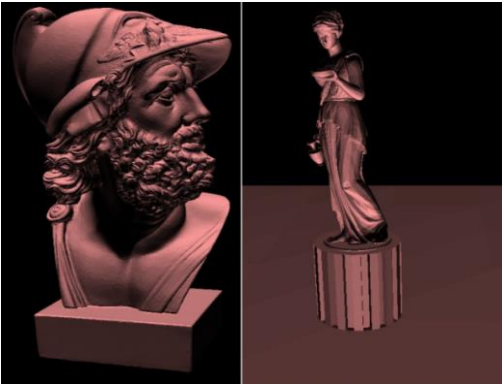


Figure 6
Part of the sample models: head (196.602 triangles), statue (95.346 triangles)

Table 1
Benchmark results

Test cases	Benchmarks (FPS)				
	Half-space algorithm variants				
	Simple rasterizer	Incremental rasterizer	Block-based rasterizer	Adaptive model	Adaptive and Bisector algorithm
C 1	112	294	456	506	517
C 2	58	165	552	552	564
C 3	60	67	51	67	69
C 4	23	37	35	40	45
C 5	69	97	106	125	132

These tests used high poly models built up from small triangles, which cannot be managed effectively by the block-based algorithm (the number of partially covered blocks increase). The adaptive model variant proved its efficiency in all cases. The logic of its triangle-orientation-based decision was able to maintain the high level performance. The fastest solution is achieved by the combination of adaptive and bisector algorithms. This made it also possible to take advantage of the reduction of empty blocks.

Conclusion

Although present day rasterization is almost exclusively performed by GPUs, we cannot forget the opportunities offered by modern CPUs. It should be recognized that certain functions can and should also be shared between GPU and CPU in order to make a more effective and robust rasterization model. To bring the two sides more closely together, this paper highlighted the basic problems of visualization. We can see that developing a fast and effective rendering model is not trivial, there are many difficulties. The authors presented some new variants of

the classic half-space triangle rasterization model, which fits much better with modern CPUs and can be a good basis for developing a more complex rasterizer, for example a hybrid pipeline between GPU and CPU. In the future, it is expected that many applications will be released using a similar technology.

Acknowledgements

This research was carried out as part of the TAMOP-4.2.1.B-10/2/KONV-2010-0001 project with support by the European Union, co-financed by the European Social Fund.

References

- [1] Bethel, Z.: A Modern Approach to Software Rasterization. University Workshop, Taylor University, December 14, 2011
- [2] TransGaming Inc: Why the Future of 3D Graphics is in Software, White Paper: Swiftshader technology, Jan 29, 2013
- [3] Microsoft Corporation: Windows Advanced Rasterization Platform (WARP) guide. 2012
- [4] Abrash, M.: Rasterization on larrabee. Dr. Dobbs Portal, 2009
- [5] Seiler, L., Carmean, D., Sprangle, E., Forsyth, T., Abrash, M., Dubey, P., Junkins, S., Lake, A., Sugerman, J., Cavin, R., Espasa, R., Grochowski, E., Juan, T., Hanrahan, P.: Larrabee: a Many-Core x86 Architecture for Visual Computing. ACM Transactions on Graphics (TOG) - Proceedings of ACM SIGGRAPH 2008 Volume 27 Issue 3, August 2008
- [6] Laine, S., Karras T.: High-Performance Software Rasterization on GPUs. High Performance Graphics, Vancouver, Canada, August 5, 2011
- [7] Akenine-möller, T., haines, E.: Real-Time Rendering, A. K. Peters. 3rd Edition, 2008
- [8] Sugerman, J., Fatahalian, K., Boulos, S., Akeley, K., and Hanrahan, P.: Gramps: A Programming Model for Graphics Pipelines. ACM Trans. Graph. 28, 4:1–4:11, 2009
- [9] Fang, L., Mengcheng H., Xuehui L., Enhua W.: FreePipe: A Programmable, Parallel Rendering Architecture for Efficient Multi-Fragment Effects. In Proceedings of ACM SIGGRAPH Symposium on Interactive 3D Graphics and Games, 2010
- [10] Agner, F.: Optimizing Software in C++ An Optimization Guide for Windows, Linux and Mac platforms. Study at Copenhagen University College of Engineering, 2014
- [11] Swenney, T.: The End of the GPU Roadmap. Proceedings of the Conference on High Performance Graphics, 2009, pp. 45-52

- [12] Coffin, C.: SPU-based Deferred Shading for Battlefield 3 on Playstation 3. Game Developer Conference Presentation, March 8, 2011
- [13] Xuzhi W., Feng G., Mengyao, Z.: A More Efficient Triangle Rasterization Algorithm Implemented in FPGA, Audio, Language and Image Processing (ICALIP), July 16-18, 2012
- [14] McCormack J., McNamara R.: Tiled Polygon Traversal Using Half-Plane Edge Functions, HWWS '00 Proceedings of the ACM SIGGRAPH/EUROGRAPHICS workshop on Graphics hardware, 2000, pp. 15-21
- [15] Chih-Hao, S., You-Ming, T., Ka-Hang, L., Shao-Yi, C.: Universal Rasterizer with Edge Equations and Tile-Scan Triangle Traversal Algorithm for Graphics Processing Units, In proceeding of: Proceedings of the 2009 IEEE International Conference on Multimedia and Expo, 2009
- [16] Valient, M.: Practical Occlusion Culling in Killzone 3, Siggraph 2011, Vancouver, Oct. 14, 2011
- [17] Olano, M., Trey, G.: Triangle Scan Conversion Using 2D Homogeneous Coordinates, Proceedings of the 1997 SIGGRAPH/Eurographics Workshop on Graphics Hardware, ACM SIGGRAPH, New York, August 2-4, 1997
- [18] Chandrasekaran, C.: Software Occlusion Culling, Intel Developer Zone, Jan 14, 2013
- [19] Leone, M., Barbagallo, L.: Implementing Software Occlusion Culling for Real-Time Applications, XVIII Argentine Congress on Computer Sciences, Oct. 9., 2012
- [20] Hill, F. S. Jr.: The Pleasures of 'Perp Dot' Products. Chapter II.5 in Graphics Gems IV (Ed. P. S. Heckbert) San Diego: Academic Press, 1994, pp. 138-148
- [21] Mileff, P., Dudra, J.: Advanced 2D Rasterization on Modern CPUs, Applied Information Science, Engineering and Technology: Selected Topics from the Field of Production Information Engineering and IT for Manufacturing: Theory and Practice, Series: Topics in Intelligent Engineering and Informatics, Vol. 7, Chapter 5, Springer International publishing, 2014, pp. 63-79
- [22] Royer, P., Ituero, P., Lopez-Vallejo, M., Barrio, Carlos A. L.: Implementation Tradeoffs of Triangle Traversal Algorithms for Graphics Processing, Design of Circuits and Integrated Systems (DCIS), Madrid, Spain; November 26-28, 2014

A Role-based Access Control Model Supporting Regional Division in Smart Grid System

Daniela Rosic, Imre Lendak, Srdjan Vukmirovic

Faculty of technical sciences, University of Novi Sad

Trg Dositeja Obradovica 6, 21000 Novi Sad, Serbia

E-mail: drosic@uns.ac.rs; lendak@uns.ac.rs; srdjanvu@uns.ac.rs

Abstract: Smart grids are modern electric power infrastructures, which incorporate elements of traditional power systems and information and communication technology (ICT), with the aim to improve the reliability, efficiency and safety requirements of critical infrastructure systems. Due to its reliance on ICT, the Smart Grid exposes electrical power systems to new vulnerabilities and security issues. Therefore, security is becoming an ever increasing concern, in the physical and ICT domain as well. Access controls are one of the most important aspects of information security and a vital element of a layered security strategy. The role-based access control (RBAC) model is widely used in complex enterprise systems which are characterized by many participants accessing the system, but with different levels of access rights depending on their specific duties and responsibilities. The existing security models, which are primarily role-based, are usually not tailored for critical infrastructure systems with specialized features, such as high numbers of equipment and devices dispersed over vast geographical regions. In order to meet the security requirements of smart grids, it is important to manage their assets on a fine level of granularity. This paper proposes an access control management system for smart grids by considering the regional division of critical assets and concept of areas of responsibility (AOR). To this end, the standardized RBAC model was extended with the aim to improve the existing access control policy with greater level of granularity from the aspect of managing electrical utilities. In this paper, we propose the RBACAOR model, which was developed and tested on the Windows operating system platform using .NET Framework role-based security, with the use of different data stores for the RBACAOR configuration, namely Active Directory (AD), AD Lightweight Directory Services (AD LDS) and Microsoft SQL Server.

Keywords: smart grid; role-based access control; regional division; area of responsibility;

1 Introduction

Smart grids are modern electric power infrastructures which incorporate elements of traditional electric power grids and information and communication technology (ICT), with the aim to improve the reliability, efficiency and safety as crucial

requirements of critical infrastructure systems. As complex systems of systems, which are integrated and interconnected over the communication network, smart grids brought along substantial benefits relating to the automation, supervision and real-time monitoring and control throughout the electric power grid, modern communications infrastructure and modern energy management techniques [1][2].

In order to provide greater reliability, efficiency and effectiveness of operations, the Smart Grid relies on a complex information and communication infrastructure to establish interconnections between different smart grid components. This complexity and the numerous interconnections (e.g. utility field crew for accessing critical operations data, control center inter-connections, etc.), make the Smart Grid a prime target of cyberattacks. Deliberate attacks are not the only threats to modern critical infrastructures (smart grids included), which might jeopardize the reliable and safe operation of these complex systems [3][4], illegitimate access and malicious attacks against smart grids can also cause data latency or data loss, thus negatively affecting operation capabilities, e.g. decision making and timely and correct responses to system events. The privacy of users might also be compromised as consumer data is now stored in business systems, which are often accessible from the Internet [5].

Because of the above listed threats, security is becoming a growing concern in smart grids, and the protection of the system and its resources from unauthorized access is of critical importance. Strict access control systems are a vital element of a layered security strategy, especially in mission-critical systems and services, where they represent a significant step towards eliminating single points of failure.

This paper proposes the $RBAC_{AOR}$ access control management system, which addresses the specific requirements of smart grids. The existing role-based access control (RBAC) model [6] will be extended to support features and security requirements specific to the smart grid environment. These specific features and requirements are analyzed in Section 2. The extended $RBAC_{AOR}$ model is discussed in Section 3. The proposed $RBAC_{AOR}$ model is applied in a smart grid environment in Section 4.

2 Security Requirements and Features of Smart Grid

2.1 Security Requirements of Smart Grids

A conventional power grid is composed of dedicated power devices which form isolated networks, with reliable and predictable communication links. In contrast, smart grids expose electrical power systems to new vulnerabilities and security issues through the introduction of complex communication networks and information systems. In order to effectively protect smart grids from cyberattacks,

strong and robust security controls are needed. Access control is one of the most important aspects of information security, critical in preserving the confidentiality, integrity and availability of information [7]. Availability might be the most important security objective in critical infrastructure systems, as a measure towards ensuring continuous, uninterrupted, real-time monitoring and control. While information integrity is an increasingly critical requirement, confidentiality was historically the least critical for electric power grids. The latest trends show, that it is becoming more significant as personal information (such as customer energy consumption data) might be available on networks and disclosure of sensitive data can lead to serious concerns regarding user privacy.

In order to meet security challenges, the National Institute of Standards and Technology (NIST) specified the following security requirements for smart grids [7]:

- Availability refers to ensuring timely and reliable access to and use of system and data. Malicious attacks targeting availability can be considered as Denial-of-Service (DoS) attacks, which intend to delay, block, or even corrupt the communication channels in the system.
- Integrity refers to guarding against information modification or destruction by unauthorized users or systems. Malicious attacks targeting the integrity of a smart grid attempt to manipulate or corrupt critical data, such as meter readings, billing information, or control commands, thus leading to the ability to negatively impact operations or even service disruption and system instability.
- Confidentiality, refers to protecting sensitive data from unauthorized access. Although malicious attacks targeting confidentiality have negligible effects on the operation of the system, serious privacy issues arise from a disclosure of customer personal information and may lead to a variety of severe consequences in the Smart Grid.

In order to address these challenging security issues in a Smart Grid, various cyber security controls are specified by the NIST. Access controls ensure confidentiality, integrity and availability of system and data by employing identification, authentication and authorization mechanisms as follows:

- Identification & Authentication: the information system(s) incorporated into smart grids have to uniquely identify and authenticate all participants (users, devices, systems) requiring access to the system.
- Authorization: the information system(s) incorporated into smart grids have to enforce authorization checks for valid users, who request access to resources within the system or between interconnected systems. Authorization may be achieved by various mechanisms to meet the needs of businesses. Different access control models are discussed in the following section.

2.2 Access Control Model Overview

The existing access control models can be classified into three broad categories: mandatory access control (MAC), discretionary access control (DAC) and role-based access control (RBAC). The MAC policy relies on security labels which have to be assigned to all users and resources in the system by system administrators, making the MAC inflexible for large systems. On the other hand, the DAC policy is based on object ownership, in which resource owners define the access privileges associated with their resources. Access control lists (ACL) are a commonly used DAC policy, enforcing specific entities attached to resources to define resource access to objects on a per user basis. Decentralized administration, as well as difficult and costly maintenance of large numbers of ACLs make the DAC less suitable for large systems.

The RBAC model was introduced as an alternative to MAC and DAC, as a model which meets the security requirements of complex enterprise systems with many participants. In RBAC, access decisions are based on the user's privileges obtained through the roles the user is authorized for. It offers easier access management, and reduces complexity and the cost of administration. These positive characteristics make RBAC suitable for systems with large numbers of users [8].

2.3 RBAC96 Model

The RBAC96 [9] is a family of role-based access control (RBAC) models which defines the scope of RBAC features included in the NIST standard. The Core RBAC (RBAC₀) defines a minimum collection of RBAC entities and entity relationships which have to be supported by any RBAC compliant system. Core RBAC elements and relations are listed below (RBAC abbreviations are enclosed in brackets):

- Users (*USERS*) are subjects (a person, computer, network) which directly interact with the system,
- Objects (*OBS*) are physical or information assets, i.e. any system resource which needs to be protected,
- Operations (*OPS*) are active processes or actions invoked by a user,
- A permission (*PRMS*) is an operation defined on an object – access right (privilege) for a resource in the system,
- A role (*ROLES*) is a job function or responsibility in the context of the system,
- User Assignments (*UA*) define user-role relations. Users are assigned to roles according to their responsibilities. Each user can be assigned one or many roles, and each role can be assigned to one or many users.

- Permission Assignments (*PA*) define permission-role relations. A role is a collection of permissions and each permission can be assigned to one or many roles. Permissions are not assigned to users directly. Instead, users obtain permissions implicitly through their roles.

The set of entities and static relations in the $RBAC_0$ model is represented in Fig. 1. The entity relationship diagram implies that a user can have one or many roles. A role is a collection of one or many permissions. Every permission is determined as a privilege needed to perform a certain operation on an object.

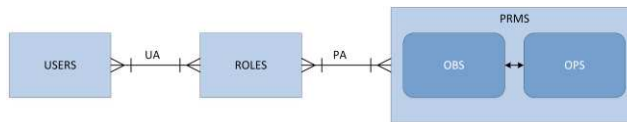


Figure 1
RBAC₀ Entity Relationships Diagram

2.4 Authorization Challenges in the Smart Grid

The Smart Grid is a collection of information systems with different operational and functional requirements [10]. The operational system represents a real-time environment for monitoring and control of Industrial Control Systems (ICS), such as a Supervisory Control and Data Acquisition (SCADA) or Distributed Control System (DCS). Beside SCADA/DCS, smart grids usually contain enterprise systems as well, which are intended for business and engineering operations, as well as to provide access from the Internet using common functions like e-mail and web services. The SCADA/DCS and the various enterprise systems are often interconnected, thereby allowing data exchange, as well as efficient operational and maintenance activities. Therefore, strict access controls are crucial for a reliable and secure integration of ICSs in corporate business processes, ensuring that all users are restricted only to the functionalities needed to accomplish their duties, as well as to disable illegitimate users from accessing the system.

Smart grids are very complex interconnected systems. For example, statistics [1] showed that there are over 2000 power distribution substations, about 5600 distributed energy facilities, and more than 130 million customers all over the US. From this it follows that smart grids are characterized by large numbers of users, critical assets and functionalities which need to be properly managed, controlled and made available to appropriate users and applications across the system. Hundreds of thousands of pieces of equipment and devices are deployed over a very large geographical area and the separation of duties and responsibilities with respect to the electricity network (e.g. by voltage level, by substation, by feeder, or by device) significantly simplifies the management of these systems.

Furthermore, sharing responsibilities among users who are assigned the same role reduces a likelihood of (configuration) errors in the system.

The RBAC96 is a rather generic access control model and does not fully meet all the security requirements of critical infrastructure systems, such as the separation of users' duties and responsibilities according to regional division of critical assets. To this end, the notion of an area of responsibility (AOR) is introduced as another level of access control in the Smart Grid environment [11]. The AOR refers to a collection of electric power system resources in a common geographic area, usually managed together. Depending on the assigned AORs, a user can be allowed to monitor, control and/or modify only certain parts of the system. Fig. 2 illustrates the relationships between electric power system resources, users and AORs. The resources belong to AORs as geographical areas, which are comprised of one or more logical areas. Users are never assigned to geographical areas directly. Instead, users are assigned to AORs as logical areas, thus obtaining a certain level of responsibility for geographical areas associated with these logical areas. In doing so, a finer granularity of access rights is allowed for users belonging to the same role. Hereafter, the notion of AOR configuration refers to both geographical areas and logical areas, as well as the relationships between them.

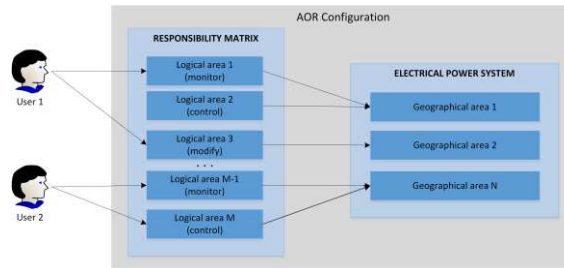


Figure 2

Area of Responsibility Configuration

3 The Formal Definition of RBACAOR Model

This sections contains a formal definition of $RBAC_{AOR}$, which extends the RBAC model and thereby meets the security requirements of smart grids discussed in Section 2. The model changes put forward in this paper propose to extend the basic set of $RBAC_0$ entities and static relations to incorporate the concept of AORs. Only the static (i.e. configuration time) entity relationships are considered, i.e. any aspect of dynamic (i.e. on-the-fly) AOR assignment is outside our scope.

3.1 The Extended Set of RBACAOR Entities and Relations

$RBAC_{AOR}$ extends the basic set of $RBAC_0$ entities with the concept of AORs [11]. The AOR entity is introduced to establish another level of access control with respect to the regional division of smart grid assets. The entity relationship diagram of the $RBAC_{AOR}$ model is represented in Figure 3. The $RBAC_{AOR}$ specific entities and relationships (marked in red) are summarized below:

- User-AOR Assignments (*UAA*) are introduced to define relationships between users and AORs. A user can be assigned zero, one or many AORs. On the other hand, an AOR can be assigned to one or many users. An AOR should not remain unassigned, as in that case a part of the Smart Grid would be unsupervised.
- Object-AOR Assignments (*OAA*) are introduced to define relationships between the smart grid assets (OBS) and AORs in the $RBAC_{AOR}$ model. Depending on the structure of the electric power system, each asset can belong to one or many AORs. One example of one-to-many relation is for assets (e.g. switchgear) placed on (or near) the boundary between the high-voltage (HV) energy management and medium-voltage (MV) distribution management system. For OBSs which are assigned to zero AORs, AOR level of access control is not considered. Similarly, an AOR can be associated with an arbitrary set of objects, depending on the business logic.

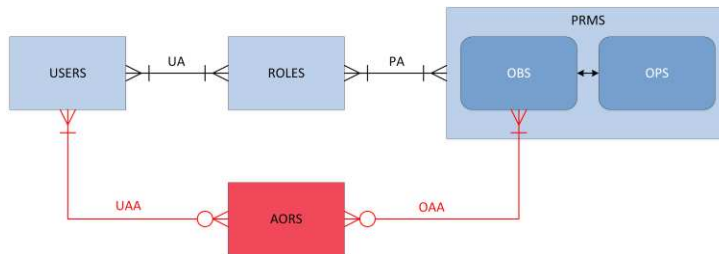


Figure 3

$RBAC_{AOR}$ Entity Relationships Diagram

3.2 AOR Responsibility Matrix for the Operation of the Smart Grid

To support different levels of responsibilities for the operation of electric power systems, the AOR entity is extended with the property *LevelOfResponsibility* which can take any of the following values: MONITORING, CONTROL and CONFIGURATION. AOR levels of responsibility are closely related to different

types of operations in the Smart Grid, which can be grouped into the below listed three categories [12][13]:

- **MONITORING** – Smart grid monitoring (such as equipment status, power flow, outage information, etc.) delivers situational awareness about power system components and performance in near real time, based on which potential problems might be identified (overloads, high/low voltage conditions, fault and outage locations) and premature equipment failures might be predicted.
- **CONTROL** – Smart grid control (i.e. operations) involves response to situations needing immediate actions in a timely and correct manner, e.g. the control of substations, transformers or feeders, configuration of different operational parameters which influence the behavior of the Smart Grid, managing incidents, etc. Control activities are usually carried out by human operators who analyze data (e.g. in the form of events and alarms in different systems) and decide which protective, preventive or corrective action should be taken.
- **CONFIGURATION** – Smart grid configuration covers activities related to modifying feeders, transformers, and other components of electric power systems, and their connectivity in the network model. Importing feeders from a Geographical Information System (GIS) might be regarded as a configuration activity in the Smart Grid.

The AOR responsibility matrix, namely the AOR Policy, is introduced to define the required level of responsibility for different types of smart grid operations. The AOR Policy is determined by the following rule: a user is allowed to execute an operation on an OBS (i.e. smart grid asset belonging to a particular AOR) only if the category of the requested operation corresponds to the AOR level of responsibility the user is assigned to. Otherwise, the user is restricted to read-only access to the requested smart grid asset.

Introducing different categories of smart grid operations does not require any additional changes to the Operation entity defined in $RBAC_0$. Instead, requiring a certain level of responsibility for the operation is accomplished by employing the AOR's property *LevelOfResponsibility* based on the (logical) category of every smart grid operation.

3.3 Class Diagram of $RBAC_{AOR}$ Authorization Framework

The set of entities and static relations in the $RBAC_{AOR}$ model is represented in Figure 4. The $RBAC_{AOR}$ class diagram is given in the context of the Microsoft Windows operating systems. A security principal is any entity which can be authenticated by the system, e.g. users and user groups. A security identifier (commonly abbreviated SID) is used as a unique, immutable identifier of users, roles, permissions and AOR entities.

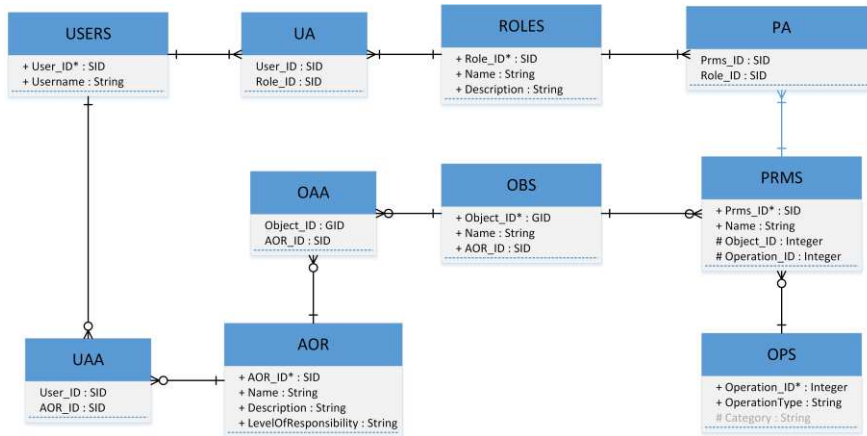


Figure 4
RBAC_{AOR} Class Diagram

4 RBAC_{AOR} Implementation

The proposed RBAC_{AOR} access control management system was implemented with Microsoft's Identity and Access Management (IAM) system and the .NET programming framework. The core component of Microsoft's IAM system is the Active Directory Domain Services (AD DS) which provides a centralized and secure identity storage and services for authentication and authorization, compliant with the Lightweight Directory Access Protocol (LDAP) protocol. Kerberos is a trusted third party authentication protocol natively included in the Active Directory (AD) environment, it ensures mutual authentication of principals and provides them with a shared session key (symmetric cipher key) to protect the confidentiality and integrity of communication channels.

The authorization mechanism ensures control of privileged operations based on information contained in a centralized identity storage. For applications running within the enterprise boundary (and having access to the AD over LDAP) the .NET authorization framework enables role-based access control mechanism in which roles are defined as AD security groups representing the relationship between users and access rights to resources in the system. Our initial analyses showed that using AD for storing the entire RBAC_{AOR} configuration would incur significant performance issues. Furthermore, AD schema changes, required to support new RBAC_{AOR} entities, would result in a significant and costly administrative burden. Furthermore, errors in AD schema changes might result in data loss or corruption. Therefore, we suggest storing the RBAC_{AOR} configuration in different types of data stores, depending on the type of data and how applications use them.

A single smart grid system was modeled in a single AD domain with its own identity storage, thus ensuring separation of users' responsibilities between different systems. However, a controlled communication must be established between systems to accomplish different business needs as discussed in Section 2. A secure integration of multiple systems occurs through the AD trust mechanism which establishes a trust relationship between identity stores, allowing for users in one domain to access resources in another domain. The remainder of this section describes details regarding the RBAC_{AOR} configuration data storage and the RBAC_{AOR} access control system.

4.1 RBAC_{AOR} Configuration Storages

Two different types of databases were used for storing RBAC_{AOR} configuration. A directory database was used for storing relatively static data which needs to be distributed among applications in a single system or between systems. A relational database was used for storing frequently changing data on a per-application basis. Figure 5 shows how the different databases were used for storing the RBAC_{AOR} entities.

Active Directory (AD) was used as a directory database to provide a centralized storage of system-specific information, such as users, roles and user-role membership, which needs to be used across the enterprise and between systems. However, AD could provide more complex access controls based on other data, including user attributes, time, data or other environmental attributes, thus allowing for a higher degree of control and more flexibility to meet the specific needs of enterprise systems.

Application-specific data (permissions, AORs and group memberships related to RBAC_{AOR}) were stored in the Active Directory Lightweight Directory Services (AD LDS). The AD LDS is a service-architected implementation of AD which runs on every machine in the system, thus alleviating run-time loads on AD and reducing network traffic. RBAC_{AOR} application data require directory schema extensions to include definitions of custom objects and attributes and storing application-specific data in AD LDS allows for flexible access control mechanism which can be extended with new application-specific entities without incurring significant risks, costs, or burdens.

Electric power system resources change more frequently compared to other RBAC_{AOR} entities (users, roles, permissions, AORs). These changes are related to creating, editing and deleting electric network elements, their connectivity and AOR memberships when the smart grid's network model is configured or integrated with other IT systems. Accordingly, Microsoft's SQL Server was used as a relational database, to store information about OBS entities, as well as, relationships between OBS and AORs.

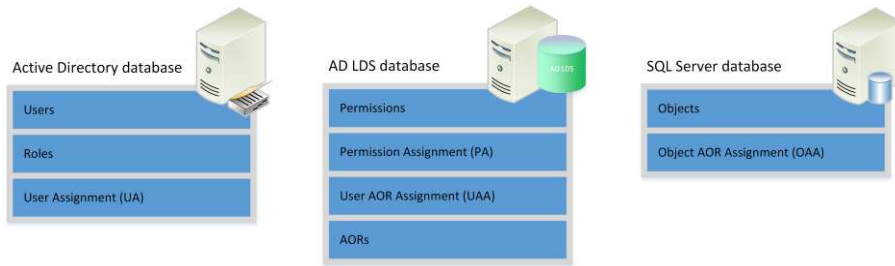


Figure 5
RBAC_{AOR} Data Storages

4.2 RBAC_{AOR} Access Control System

As presented in Figure 6, the RBAC_{AOR} access control system is comprised of two processes, authentication and authorization, which are combined to ensure that resources are accessed only by authorized users.

The RBAC_{AOR} authentication framework utilizes .NET Integrated Windows authentication (IWA) as a first step toward gaining access to the system. IWA is based on the Secure Protocol Negotiation (SPNEGO) security package embedded in the Windows operating systems. Although the SPNEGO can interface with both the Kerberos and the NTLM authentication protocols, Kerberos was chosen as the best suited authentication protocol for intranet environments where both clients and servers are in the same domain or trusted domains. Kerberos is a widely-adopted network authentication protocol, in which client and server mutually authenticate each other based on a reliable third-party. Kerberos ensures session confidentiality and integrity by using session keys [14].

The RBAC_{AOR} security principal is a result of successful authentication and consolidates identity information from multiple data storages (AD and AD LDS) which are combined in real-time. The RBAC_{AOR} authorization framework is based on .NET Framework role-based security which has been adapted to specific data stores of system and applications, such as AD, AD LDS and SQL Server. The RBAC_{AOR} authorization framework is comprised of two independent authorization processes which are combined to determine the final access control decision based on information encapsulated into the RBAC_{AOR} security principal:

1. The Role-Based Access Control flow is executed by a user request to execute operation Op on an OBS. A RBAC access decision is determined by checking whether the RBAC_{AOR} security principal has a permission which defines the privilege for the required resource.
2. The AOR access control flow is executed by a user request to access OBS, belonging to an area of responsibility AOR. An AOR access decision is determined by checking whether the RBAC_{AOR} security

principal is a member of at least one logical area associated with the AOR to which the object (OBS) belongs.

The final access control decision is determined based on the RBAC access decision and AOR access decision so that access is authorized only when both RBAC and AOR access decisions are determined as allowed. For authorized users, the level of responsibility is further considered based on the AOR Policy. As explained in Section 3, a user is allowed to execute an operation Op on an object (OBS) only if the user is assigned to at least one logical area (related to the AOR to which object belongs) with the level of responsibility which corresponds to the category of the requested operation. Otherwise, the user is restricted to read-only access to the requested operation on the smart grid asset.

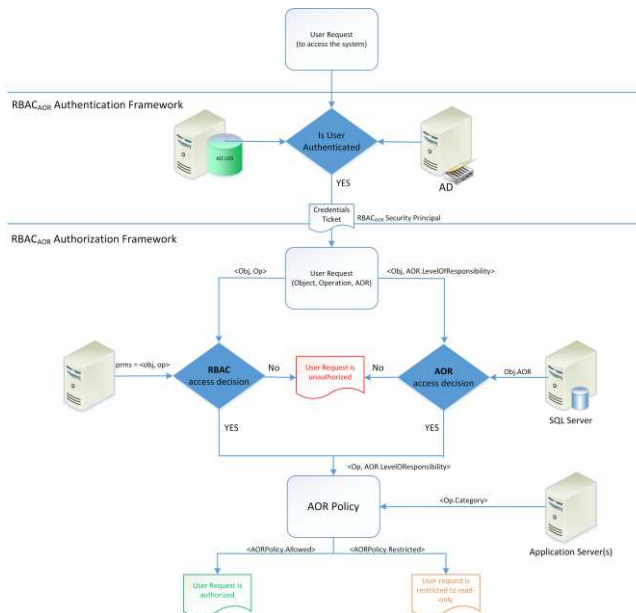


Figure 6
RBAC_{AOR} Access Control System

Conclusions

This paper addresses the problem of access control in smart grids, as one of the most important aspects in preserving the confidentiality, integrity and availability of smart grid (and critical infrastructure systems in general) assets. First, we analyzed features and security requirements specific to smart grids, in order to define the strengths and weaknesses of the existing access control models, with emphasis on the RBAC model, as the most commonly used model in large enterprise systems. While analyzing the security requirements of smart grids, we deduced that the existing access control model based on a user's roles and

responsibilities within the system does not cover every security requirement of modern electric power grids with large numbers of equipment and heterogeneous devices dispersed over vast geographical regions. Therefore, this paper presents an extension to the standard $RBAC_0$ model with the concept of the area of responsibility (AOR). The proposed $RBAC_{AOR}$ model is comprised of two separate components: role-based access control flow and AOR-based access control flow. Both components check independently whether users have appropriate access rights, and make the final access decision based on the combination of access decisions made by each subcomponent.

The $RBAC_{AOR}$ model extends the currently available role-based access controls, to provide an efficient and highly secure access control method designed specifically for smart grids. The proposed access control model was implemented with Microsoft technologies and can be easily integrated into existing role-based access control systems which are based on Active Directory security services.

Hierarchical RBAC or constrained RBAC were outside the scope of this research, i.e. only the Core RBAC features and the set of static relationships were taken into consideration. Making access control decisions based on user and/or session attributes to enforce different constraints depending on the context in which limitations are imposed was (also) outside of the scope of this paper. Some examples include roles and AORs which can be activated/deactivated for a given user, depending on the workstation the user has logged into, or dynamic AOR assignment to support transfer of duties during regular system activities. Activities related to real time simulation in smart grids in order to analyze the behavior of a network which evolves might also require specific authorization policies. The authors intend to explore these questions as part of their future research. It is important to note that the $RBAC_{AOR}$ model proposed in this paper is flexible and extensible, and the authors are confident that it will easily incorporate solutions for these additional requirements.

References

- [1] Xu Li, Xiaohui Liang, RongXing Lu, Xuemin Shen, Xiaodong Lin, Haojin Zhu: "Securing Smart Grid: Cyber Attacks, Countermeasures, and Challenges", IEEE Communications Magazine, August 2012, pp. 38-45
- [2] Wenye Wang, Zhuo Lu: "Survey Cyber Security in the Smart Grid: Survey and Challenges", Computer Networks: The International Journal of Computer and Telecommunications Networking, April 2013, pp. 1344-1371
- [3] H. Melvin: "A Role of ICT in Evolving SmartGrids", 10th International Conference on Digital Technologies (DT), July 2014, pp. 235-237
- [4] Ruofei Ma, Hsiao-Hwa Chen, Yu-Ren Huang, Weixiao Menart: "Smart Grid Communication: Its Challenges and Opportunities", IEEE Transactions on Smart Grid, February 2013, pp. 36-46

- [5] Zhang Zhigang, Liu Hao, Niu Shuangxia, Mo Jiansong: “Information Security Requirements and Challenges in Smart Grid”, 6th IEEE Joint International Information Technology and Artificial Intelligence Conference (ITAIC), August 2011, pp. 90-92
- [6] National Institute of Standards and Technology (NIST): “Role-based Access Control”, 2003
- [7] National Institute of Standards and Technology (NIST), NISTIR 7628: “Guidelines for Smart Grid Cyber Security: Vol. 1, Smart Grid Cyber Security Strategy, Architecture, and High-Level Requirements”, 2010
- [8] Maria B. Line, Inger Anne Tondel, Martin G. Jaatun: “Cyber Security Challenges in Smart Grids”, 2nd IEEE PES International Conference and Exhibition on Innovative Smart Grid Technologies (ISGT Europe), December 2011, pp. 1-8
- [9] David F. Ferraiolo, Ravi Sandhu, Serban Gavrila, D. Richard Kuhn, Ramaswamy Chandramouli: “Proposed NIST Standard for Role-based Access Control”, ACM Transactions on Information and System Security, August 2001, pp. 224-274
- [10] J. Zerbst, M. Schaefer, I. Rinta-Jouppi: “Zone Principles as Cyber Security Architecture Element for Smart Grids”, IEEE PES Innovative Smart Grid Technologies Conference Europe (ISGT Europe), October 2010, pp. 1-8
- [11] D. Rosic, U. Novak, S. Vukmirovic: “Role-based Access Control Model Supporting Regional Division in Smart Grid System”, 5th International Conference on Computational Intelligence, Communication Systems and Networks (CICSyN), June 2013, pp. 197-201
- [12] A. P. S. Meliopoulos, G. J. Cokkinides, Renke Huang, E. Polymeneas, P. Myrda: “Grid Modernization: Seamless Integration of Protection, Optimization and Control”, 47th Hawaii International Conference on System Sciences (HICSS), January 2014, pp. 2463-2474
- [13] Applied Communication Sciences: “Smart Grid – Leveraging Distributed Operations Capabilities”, 2012
- [14] Dahui Hu, Zhiguo Du: “An Improved Kerberos Protocol Based on Fast RSA Algorithm”, IEEE International Conference on Information Theory and Information Security (ICITIS), December 2010, pp. 274-278