

EDITORIAL BOARD

Editor-in-Chief

Igor B. Shubinsky, PhD, D.Sc in Engineering, Professor, Expert of the Research Board under the Security Council of the Russian Federation, Deputy Head of Integrated Research and Development Unit, JSC NIIAS (Moscow, Russia)

Deputy Editor-in-Chief

Schäbe Hendrik, Dr. rer. nat. habil., Chief Expert on Reliability, Operational Availability, Maintainability and Safety, TÜV Rheinland InterTraffic (Cologne, Germany)

Deputy Editor-in-Chief

Mikhail A. Yastrebenetsky, PhD, D.Sc in Engineering, Professor, Head of Department, State Scientific and Technical Center for Nuclear and Radiation Safety, National Academy of Sciences of Ukraine (Kharkiv, Ukraine)

Deputy Editor-in-Chief

Way Kuo, President and University Distinguished Professor, Professor of Electrical Engineering, Data Science, Nuclear Engineering City University of Hong Kong, He is a Member of US National Academy of Engineering (Hong Kong, China) (Scopus) (ORCID)

Executive Editor

Aleksey M. Zamyshliaev, PhD, D.Sc in Engineering, Deputy Director General, JSC NIIAS (Moscow, Russia)

Technical Editor

Evgeny O. Novozhilov, PhD, Head of System Analysis Department, JSC NIIAS (Moscow, Russia)

Chairman of Editorial Board

Igor N. Rozenberg, PhD, Professor, Chief Research Officer, JSC NIIAS (Moscow, Russia)

Cochairman of Editorial Board

Nikolay A. Makhutov, PhD, D.Sc in Engineering, Professor, corresponding member of the Russian Academy of Sciences, Chief Researcher, Mechanical Engineering Research Institute of the Russian Academy of Sciences, Chairman of the Working Group under the President of RAS on Risk Analysis and Safety (Moscow, Russia)

EDITORIAL COUNCIL

Zoran Ž. Avramovic, PhD, Professor, Faculty of Transport and Traffic Engineering, University of Belgrade (Belgrade, Serbia)

Leonid A. Baranov, PhD, D.Sc in Engineering, Professor, Head of Information Management and Security Department, Russian University of Transport (MIIT) (Moscow, Russia)

Alexander V. Bochkov, D.Sc in Engineering, Deputy Head of Integrated Research and Development Unit, JSC NIIAS (Moscow, Russia)

Konstantin A. Bochkov, D.Sc in Engineering, Professor, Chief Research Officer and Head of Technology Safety and EMC Research Laboratory, Belarusian State University of Transport (Gomel, Belarus)

Boyan Dimitrov, Ph.D., Dr. of Math. Sci., Professor of Probability and Statistics, Kettering University Flint, (MICHIGAN, USA) (ORCID)

Valentin A. Gapanovich, PhD, President, Association of Railway Technology Manufacturers (Moscow, Russia)

Victor A. Kashtanov, PhD, M.Sc (Physics and Mathematics), Professor of Moscow Institute of Applied Mathematics, National Research University "Higher School of Economics" (Moscow, Russia)

Sergey M. Klimov, PhD, D.Sc in Engineering, Professor, Head of Department, 4th Central Research and Design Institute of the Ministry of Defence of Russia (Moscow, Russia)

Yury N. Kofanov, PhD, D.Sc. in Engineering, Professor of Moscow Institute of Electronics and Mathematics, National Research University "Higher School of Economics" (Moscow, Russia)

Achyutha Krishnamoorthy, PhD, M.Sc. (Mathematics), Professor Emeritus, Department of Mathematics, University of Science and Technology (Cochin, India)

Eduard K. Letsky, PhD, D.Sc in Engineering, Professor, Head of Chair, Automated Control Systems, Russian University of Transport (MIIT) (Moscow, Russia)

Victor A. Netes, PhD, D.Sc in Engineering, Professor, Moscow Technical University of Communication and Informatics (MTUCI) (Moscow, Russia)

Ljubiša Papić, PhD, D.Sc in Engineering, Professor, Director, Research Center of Dependability and Quality Management (DQM) (Prijevor, Serbia)

Roman A. Polyak, M.Sc (Physics and Mathematics), Professor, Visiting Professor, Faculty of Mathematics, Technion – Israel Institute of Technology (Haifa, Israel)

Dr. Mangey Ram, Prof. (Dr.), Department of Mathematics; Computer Science and Engineering, Graphic Era (Deemed to be University), Главный редактор IJMEMS, (Dehradun, INDIA) (ORCID)

Boris V. Sokolov, PhD, D.Sc in Engineering, Professor, Deputy Director for Academic Affairs, Saint Petersburg Institute for Informatics and Automation of the Russian Academy of Sciences (SPIIRAS) (Saint Petersburg, Russia)

Lev V. Utkin, D.Sc in Engineering, Professor, Director, Institute of Computer Science and Technology, Peter the Great St. Petersburg Polytechnic University (Saint Petersburg, Russia)

Evgeny V. Yurkevich, PhD, D.Sc in Engineering, Professor, Chief Researcher, Laboratory of Technical Diagnostics and Fault Tolerance, ICS RAS (Moscow, Russia)

THE JOURNAL PROMOTER:

"Journal "Reliability" Ltd

*It is registered in the Russian Ministry of Press,
Broadcasting and Mass Communications.
Registration certificate ПИ 77-9782,
September, 11, 2001.*

*Official organ of the Russian Academy
of Reliability*

Publisher of the journal

LLC Journal "Dependability"

Director

Dubrovskaya A.Z.

The address: 109029, Moscow,
Str. Nizhegorodskaya, 27, Building 1, office 209
Ltd Journal "Dependability"
www.dependability.ru

Printed by OOO Otmara.net. 2/1 bldg 2,
Verkhniaya Krasnoselskaya St., floor 2, premise II,
rooms 2A, 2B, 107140, Moscow, Russia.

Circulation: 500 copies.

Printing order

Papers are reviewed. Signed print 17.09.2021,
Volume , Format 60x90/8, Paper gloss

The Journal is published quarterly since 2001.
The price of a single copy is 1045 Rubles, an annual
subscription costs 4180 Rubles.
Phone: +7 (495) 967 77 05.
E-mail: dependability@bk.ru.

Papers are reviewed.

Papers are published in author's edition.

CONTENTS

System analysis in dependability and safety

Pokhabov Yu.P. On the dependability of highly critical non-recoverable space entities with short operation life. Case study of single-use mechanical devices	3
Antonov A.V., Chepurko V.A., Cherniaev A.N. Method of estimating the size of an SPTA with a safety stock	13
Belousova M.V., Bulatov V.V. Estimation of the failure flow of a set of passenger car doors	20
Zharkov M.L., Pavidis M.M. Simulation of railway marshalling yards using the methods of the queueing theory	27

Discussion of dependability terminology

Zelentsov B.P. Use of the terms “estimate” and “definition” in dependability-related standards	35
---	----

Safety. Risk management. Theory and practice

Grachev Ya.L., Sidorenko V.G. Steganalysis of the methods of concealing information in graphic containers	39
Braband J., Shäbe H. Risk Model of German Corona Warning App – Reloaded	47
Pronevich O.B., Zaitsev M.V. Intelligent methods for improving the accuracy of prediction of rare hazardous events in railway transportation	54
In memory of A.M. Zamyshliaev	65
Gnedenko Forum	66

On the dependability of highly critical non-recoverable space entities with short operation life.

Case study of single-use mechanical devices

Yuri P. Pokhabov, Joint Stock Company NPO PM – Maloe Konstruktorskoye Buro (AO NPO PM MKB), Zheleznogorsk, Krasnoyarsk Krai, Russian Federation
pokhabov_yury@mail.ru



Yuri P. Pokhabov

Abstract. Aim. To consider matters of dependability of highly critical non-recoverable space products with short operation life, whose failures are primarily caused by design and process engineering errors, manufacturing defects in the course of single-unit or small-scale production, as well as to define the methodological approach to ensuring the required reliability. **Methods.** Options were analysed for improving the dependability of entities with short operation life using the case study of single-use mechanical devices and the statistical approaches of the modern dependability theory, special methods of dependability of actuated mechanical assemblies, FMEA, Stage-Gate and ground experiments on single workout equivalents for each type of effect. **Results.** It was concluded that additional procedures need to be conducted for the purpose of predicting, mitigation and (or) eliminating possible failures as part of the design process using exactly the same approaches that cause failures, i.e., those of design and process engineering. The engineering approaches to dependability are based on early identification of possible causes of failures, which requires a qualified and systemic analysis aimed at identifying the functionality, performance and dependability of an entity, taking into account critical output parameters and probabilistic indicators that affect the performance of the required functions with the allowable probability of failure. The solution is found using a generalized parametric model of operation and design engineering analysis of dependability. **Conclusion.** For highly critical non-recoverable space entities with short operation life, the reliability requirements should be considered primarily in terms financial, economic, safety-related and reputational risks associated with the loss of spacecraft. From a design engineer's standpoint, the number of nines after the decimal point (rounded to a smaller number of nines for increased confidence) should be seen as the indicator for the application of the appropriate approaches to ensuring the required reliability at the stage of product design. In case of two nines after the decimal point it is quite acceptable to use analytical and experimental verification techniques common to the aerospace industry, i.e., dependability calculations using the statistical methods of the modern dependability theory and performance indicators, FMEA and Stage-Gate, ground experiments on single workout equivalents for each type of effect. As the required number of nines grows, it is advisable to also use early failure prevention methods, one of which is the design engineering analysis of dependability that enables designers to adopt substantiated design solutions on the basis of engineering disciplines and design and process engineering methods of ensuring quality and dependability. The choice of either of the above dependability strategies is determined solely by the developer's awareness and understanding of potential hazards, which allows managing the risk of potential rare failures or reasonably refusing to do so.

Keywords: dependability calculation, Stage-Gate, FMEA, actuated mechanical assemblies, single-use devices, spacecraft, design engineering analysis of dependability (DEAD).

For citation: Pokhabov Yu.P. On the dependability of highly critical non-recoverable space entities with short operation life. Case study of single-use mechanical devices. Dependability 2021;3: 3-12. <https://doi.org/10.21683/1729-2646-2021-21-3-3-12>

Received on: 26.05.2021 / **Upon revision:** 15.07.2021 / **For printing:** 17.09.2021.

Introduction

In the process of insertion, the configuration of a modern spacecraft undergoes four changes of its kinematic state [1]:

1) operation as part of the launch vehicle with compactly folded structures in the launching position (the satellite is installed on the rocket and its folding structures are arranged within specified dimensions and fixed on the body);

2) separation from the launch vehicle and orbital flight with folding structures in the launch position (the satellite is decoupled and is at a safe distance from the rocket, yet its structures remain arranged and fixed on the body until the preparation to deployment is complete);

3) deployment of the structures from the launch position to service position by means of mechanical devices (the mechanical connections to the body housing are removed and the structural elements execute the required motions taking the specified cantilever position relative to the body);

4) the operation of the on-board systems and satellite equipment for the intended purpose during the specified lifetime of active existence with the open structures in the working position.

The above sequence of satellite state changes is defined by the conditions and restrictions for its delivery to Earth orbit by a multiple-stage rocket [2]. Only after all the mechanisms have operated – separated and deployed the folding structures in the specified configuration – the spacecraft is able to operate normally in orbit. Otherwise, all the efforts associated with the construction and launch of a spacecraft-carrying rocket lose their effectiveness, sometimes even their meaning.

Space-based mechanisms are non-recoverable systems, therefore the cost of failure in the process of separation from the launch vehicle and deployment of mechanical devices is a partial or complete loss of spacecraft functionality even before the start of the operation it was created and launched for [3]. As the history of space launches shows, failures of single-use mechanical devices are not very rare. For instance, the proportion of failures at the stage of satellite deployment may be as high as 10.05%, and 12.8% at the stage of separation from the launch vehicle [4]. At the same time, practically every time the satellite sustains various degrees of damage (except in cases of self-deployment after failures caused by thermal effects, for example, when Kiku 8 deployed its antennas). For example, three months of spin-up manoeuvres following the non-deployment of the C-band antenna on Anik E2 resulted in excess consumption of an amount of fuel corresponding to one year of normal in-orbit operation. The incomplete deployment of solar arrays on Telstar 14, Telstar 14R and Intelsat 19 caused power short-

ages, which entailed a forced shutdown of a part of the transmitter-receiver devices of the payload (for example, on Telstar 14R, 17 transponders out of 41 were disabled). The non-deployment of solar arrays caused the loss of the \$190-mil. Sinosat 2 and the \$250-mil. Chinasat 18 that could not commence their intended operation. Every year, in the world there are at least 1 or 2 failures of single-use mechanical devices in the course of spacecraft separation and deployment in the orbital phase of the flight, while the average probability of failure is as high as 0.004 per year [5]. At the same time, according to OST 92-4339, the reliability of the deployment and retention mechanisms is to be not less than 0.999 with the confidence level of 90%, while the specified pointwise value of probability of devices operation for modern long-operation spacecraft is 0.9995 [6-8].

The conclusions are simple and disappointing. In more than 60 years of space exploration, no scientifically substantiated methods have been developed for designing and building single-operation mechanisms with the required reliability. Additionally, even a long (over the last 20÷30 years) lack of failures of the separation and deployment mechanisms designed by certain manufacturers cannot be considered as indisputable evidence of the faultless methods of ensuring dependability due to the small samples of statistical data. In particular, given the required reliability of mechanical devices over 0.9995, the accident-free launch of up to 10 devices per year (maximum about 300 devices over 30 years), which is the standard volume of production of one of Russia's largest developers, by itself does not guarantee a reliability level of 0.9995 even with the confidence level of 15% [5, 9, 10]. That is assuming that losses of single-use mechanisms are practically unaffected by such sources of uncertainty as the ageing and degradation of materials and connections resulting from long exposure to space flight factors. In most cases, the time to failure of such products is minimal. That is the time of launching into orbit and deployment from the launch position that, in total, does not exceed dozens of minutes and does not suppose reactivation in orbit [5, 8]. Accordingly, the failures of mechanical devices are determined mainly by design and manufacturing errors, as well as non-observance of the conditions of zero-defect production of single and/or small-batch products [5, 11-14]. The prevention¹ of such failures mainly depends on the degree of substantiation and establishment of the indispensable and sufficient requirements in the

¹ Prevention of failures: The implementation – in the course of the construction (upgrade), manufacture and operation of products – of a set of managerial and technical measures enabling the prevention, detection, investigation and elimination of the causes of product failures [OST 134-1012-97, Section 4].

design documentation for the purpose of manufacture and appropriate supervision of the key values of critical elements at all life cycle stages [15].

Given the above circumstances, let us consider ways of increasing the dependability of highly critical non-recoverable products with short operation life and one of the methodological approaches to ensuring the required functional reliability of single-use mechanical devices of spacecraft.

Capabilities of the modern dependability theory

First, literally all regulatory documents and scientific and methodological literature require calculating the dependability on the basis of undependability known from experience, i.e., a posteriori knowledge on possible failures [16, 17]. It is believed that dependability calculation is to be based on the presumption of failures, that are allegedly inevitable by definition, therefore, in order to calculate the dependability, it is required to know the statistical probability of failure of the entire product or at least the actual undependability of its components and elements in the specified modes and conditions of application [18]. If there are no known dependability (in reality, undependability) indicators, then, according to the modern dependability theory, they should be produced using statistical methods [19-21]. No other way is allowed by the requirements of such standards as, for example, GOST 27.002, GOST 27.301, GOST RO 1410-001, etc.¹

Second, as regards single-use mechanical devices of spacecraft, ensuring a reliability of, e.g., 0.9995 it is required to hold at least 9995 independent tests (experiments) under uniform conditions. In other words, according to regulatory documents, it is required to deploy in orbit (not on the ground, otherwise the uniform conditions will not be observed) at least 9995 mechanical devices (not testing one device 9995 times, otherwise the test independence is not observed). All of that is only to confirm, that a single normal opening will occur with the required dependability. Let us assume a 0.9995 reliability would be sufficient with a confidence level of 0.9, but even then, the number of independent tests in uniform conditions must not be less than 4605 [22] (see example in GOST R 27.003 for identifying the minimal scope of statistical tests as part of dependability-related contracts). If we use the dependability calculation method based on known dependability indicators of components and elements, the number of required statistical tests in outer space will be considerably higher than for the mechanical devices themselves, since they consist of tens and hundreds of components in the form of the simplest

mechanisms and devices, the number of requirements to the dependability of which grows exponentially with respect to the number of functional elements that affect the overall dependability of the system [23]. It is obvious that it is almost impossible to obtain reliable data for calculating the dependability of highly vital mechanical devices using statistical methods of the modern dependability theory due to financial and economic reasons (as of 2018, the cost of launching 1 kg of freight was \$20-30 ths, as of 2020, it was \$15-17 ths [24]).

Third, the humanity simply does not possess the required numbers of equivalent items to calculate the dependability of mechanical devices with a reliability close to 0.9999. The total number of satellites successfully launched worldwide between 1958 and 2010 is 6264 [25]. Between 2011 and 2016, 1153 more spacecraft were launched [5]. Even if we ignore the requirement of sample homogeneity, we still cannot rely on the reliability of statistical data for the purpose of calculating the reliability of deployment of, for example, solar panels (installed on almost every satellite) at the level of 0.9995.

Fourth, in the aerospace industry it is conventionally believed that flight-qualified products are dependable [9]. However, the high requirements for the dependability indicators in cases of small sample sizes do not correspond to the statistical approaches of the modern dependability theory. A failure that is acceptable, for example, for 10000 tests (experiments), may occur at the time of any of the tests, and, let us suppose, if 100 successive tests went successfully, there is no guarantee the 101-th does not end in a failure. In this case, it can be said that the product has confirmed its performance 100 successive times (but in no case conclusions concerning dependability can be made). Therefore, without the scientific and methodological substantiation of the feasibility of the required reliability (i.e., without additional analysis and/or simulations confirming the performance of the required functions with no failures), from an engineer's standpoint, it would be simply careless to draw any conclusions regarding the dependability of highly vital products in cases of low operation life.

Thus, using the statistical approaches of the modern dependability theory alone is not acceptable for the purposes of ensuring high operational dependability of single-use mechanical systems. Obviously, in this case the causes, rather than the consequences (statistics) of failures must be first identified. Therefore, methods of engineering analysis and dependability calculation are required that would be based on the physical phenomena described by physical theories. That would enable the construction of mathematical models of loss (or retention) of an object's performance with the change of its internal state in the specified modes and conditions of application [26].

¹ See terms related to the methods for dependability identification [articles 3.7.9-3.7.11].

Special methods for ensuring the dependability of single-use mechanical devices

The method of dependability calculation of the mechanical parts of moving structures of spacecraft was first published in 1978-1979 [27, 28]. In addition to the durability, the dependability calculation was proposed that is based on identifying the probability of excess drive moments (forces) of actuators over the resistance moments (forces) in the path of motion of the executive devices, as well as the calculation of the overall dependability of mechanisms on the basis of the phantom item (unit) model [28]. Later, mechanism dependability was calculated taking into account the margin of drive moments (forces) [29-33] similarly to the deterministic calculations for strength based on the safety factors and safety margin [34]. Abroad, the compliance with requirements for margin of drive moments (forces) is an integral part of all standards for designing moving mechanical assemblies (MMA) intended for space application. In 1975, the standard values of the margin of drive moments (forces) were specified in the military standard MIL-A-83577, later, in the civil standards AIAA S-114-2005, NASA-STD-5017A and ECSS-E-ST-33-01C. In Russia, there are no official standards (GOST, GOST R, OST, STP, STO) for designing mechanical devices taking into account the margin of drive moments (forces), but the years-old application practice is that deployment drives are selected on the basis of the requirement of a margin of drive moments (forces) not less than 100% (2:1 ratio) of the worst value of the resistance moments (forces) at any point of the path of motion assuming zero kinetic energy [32, 33]. It is commonly believed that if the specified reliability coefficients, strength margins, drive moments (forces) and conditions of successful confirmation of the criteria of experimental optimization (as defined in GOST R 58630) are observed, the specified reliability of deployment and retention of mechanical devices is ensured by default [29-31].

However, studies of the actual causes of failure show that in the vast majority of cases they are rare in terms of their nature that is defined by an unfavourable combination of manufacturing tolerances, unaccounted factors of technological heredity, application modes and external effects [5, 15]. Such failures can be caused, for example, by sudden disappearance of gaps in kinematic pairs (Kiku 8, Soyuz TMA-17M), unfavourable combination of production factors (Intelsat 19), manufacturing defects (Kanopus-ST, Progress M-19M), foreign objects in the deployment mechanism (Skylab, Telstar 14, Telstar 14R), failures of deployment actuators (EchoStar 4), unauthorized deployment (Resurs-P no. 3), design and manufacturing errors (Mayak), cold welding (Galileo), etc. The

practice shows that dependability calculations using the statistical methods of the modern dependability theory and performance parameters (from the recommended list in OST 92-0290), as well as successful ground experiments on single workout equivalents for each type of effect, are unable to prevent the risk of rare failures [16, 35]. Modern methods of experimental optimization are not intended for identifying and emulating loading cases that correspond to critical combinations of critical states of a product, factors of modes and external effects [5]. Moreover, for small probabilities of failures (not more than 0.01), the total error of dependability evaluation based on the results of experimental optimization can be as high as an order of magnitude of the valid digit, while in terms of engineering calculations an error of not more than 5÷10% [5, 36-38] is acceptable.

The Stage-Gate concept

According to the Stage-Gate concept, the execution of any project¹ is defined by sequential execution of cross-functional actions and activities (stage) separated from each other with decision points (gate) that lead to the next stage of the work plan (in the stage-gate system) [39].

In fact, it refers to process project management standards that, unlike those adopted in Russia (GOST, GOST R, OST, STP, STO), establish the order and procedures for appropriate decisions and actions. In particular, the process principle is at the foundation of the ESA standards intended for the purpose of management, engineering and quality assurance in space projects, for instance for space mechanisms (ECSS-E-ST-33-01C).

Despite the obvious benefits of the Stage-Gate-based process standards, i.e., the availability to “average” engineers for the purpose of achieving the required quality and dependability, when all of their decisions and actions are regulated by a set of pre-defined (by someone else) procedures, a thoughtless execution of formalized instructions can lead to the loss of the physical significance of decisions and the purpose of actions. For example, in the process of development of a mechanism with any particular principle of action, process standards are certainly useful, but if the physical principles of such mechanism’s operation change, it becomes necessary to promptly compensate the shortcomings of the used procedures in the standards. This can be made possible by applying engineering techniques based on strictly defined algorithm-based procedures or by the engineers’ heuristics. In the first case, that means a quick adjustment of the existing engineering methodology, in the second case, that means a relatively long way of trial and error

¹ Project: A temporary enterprise aimed at creating a unique product, service or deliverable [PMBOK, Glossary].

involving the accumulation and generalization of the behaviour patterns of new products for the purpose of enabling the required properties [40].

The FMEA analysis

The purpose of failure mode and effects analysis (FMEA) is to enable the detection and elimination of technical problems within complex systems by examining each type of failure of any critical component. FMEA and its versions: DFMEA, PFMEA, and MFMEA are based on brainstorming or expert evaluation of the types and consequences of failures of critical elements, complemented, if required, by failure mode, effects and criticality analysis (FMECA) or failure modes, effects and diagnostic analysis (FMEDA).

The FME[C,D]A method involves the following steps:

- definition of the structure of the analysed object (structural analysis);
- identification of the possible critical event scenarios (functional analysis);
- execution of the analysis to determine the types, effects and causes of failures with risk assessment for the purpose of preventive or corrective action (FMEA), FMEA-based calculation of safety indicators, i.e., risk priority or failure criticality (FMECA), FME[C]A-based identification of the failure rate (FMEDA) for dependability calculation;
- evaluation and documentation of analysis results (FMEA, FMECA or FMEDA).

FME[C,D]A analysis is performed by a cross-functional team of domain experts (e.g. designer, engineering technologist, assembler, tester, supervisor, etc.) of up to 7 or 8 people who possess practical experience and high level of professionalism [41]. The principles of FMEA team building and work organization are defined in standards and guidelines, e.g., GOST R 51814.2, STB 1506, RD 03-418-01, etc.

FMEA analysis and its extended variants (FMECA, FMEDA) are performed by experts using formalized algorithms and procedures for obtaining subjective semi-quantitative estimates (based on consequence significance ratings, probability of occurrence and detection) of potential failures (faults). The experts normally have different professional views on the analysed object that does not always correspond to the understanding of how exactly and in what conditions such object operates [42]. Experts do not know (they do not have to know according to FMEA standards) the design concept aimed at solving specific technical problems, therefore they evaluate the consequences of defects on the basis of external features (indicators) that a consumer can notice and the experts can understand from the standpoint of personal professional qualities (knowledge, qualification and experience).

Meanwhile, the designers' intent is very closely associated with establishing and substantiating the output parameters of critical elements of an object within the permissible range of values [17, 36, 37]. Moreover, if the FMEA standards (for example, GOST R 51814.2) require defining the types of potential failures in physical and technical terms (crack, deformation, jamming, destruction, leakage, etc.), then the consequences of failures are recommended to be described in the consumer's language (what he/she can notice or experience), e.g., noise, incorrect operation, instability, intermittent operation, etc. [43]. The FMEA results are either not at all or indirectly related to the output parameters and their allowable ranges that are in one way or another defined by the designer.

The approach to ensuring the dependability of single-use mechanical devices

Given that the methods of the modern dependability theory do not enable a sufficiently accurate solution of the problems of dependability of highly vital products with short operation life due to non-applicability of statistical approaches, while special and auxiliary methods are not designed for identifying the causes and assessing the risks of rare failures, it is only left to predict, mitigate or prevent possible failures at the design stage using exactly the same approaches that cause failures, i.e., those of design and process engineering.

According to the principles of rational design, a design¹ and any of its structural elements should be considered from the standpoint of them performing strictly defined functions that were originally conceived and implemented by the designer by adopting and executing specific solutions (design, engineering, design and engineering, process engineering)². Such solutions are based on a physical understanding of the world and the use of design and engineering methods for their implementation as part of technical objects. In this case, each of the designer's decisions that are potentially capable of causing a failure must be substantiated. Each argument must be compliant with the designer's logic of reasoning that he/she understands in the context of ensuring a failure-free operation of the product.

It may be advisable to use methods of parametric modelling of products based on the available diagrams

¹ Design: A device, the mutual arrangement of parts of an object, machine, instrument defined by its purpose and involving a method of ensuring the connection, interaction of parts, as well as the material the individual parts (elements) must be made of [GOST R 57945-2017, Article 2.66].

² According to definitions of the respective terms associated with the word "decision" per GOST R 57945.

(design layout, structural, etc.), sketches, drawings, 3D models in order to confirm the solutions. In this case, any graphic, text-and-graphics or digital design models must be represented in the form of a parametric model, the modification of whose parameters enables the fulfilment by the product of all required functions.

Based on the principles of physicality (causal connections)¹ and physical necessity (consistency with the laws of nature)² it is not difficult to represent the performance of the required functions by the product on the basis of the parametric model that describes its functionality, performance and dependability [17, 36, 37, 44]. The logical sequence of reasoning is as follows. If the design is represented as a set of output parameters that characterize the performance of the required functions (i.e. the functionality), each design parameter is defined based on a combination of the modes and conditions of application (i.e., performance), while the modification of the values of the design parameters over time is restricted within the allowable range (i.e., dependability), a generalized parametric model of the product's operation can be obtained, in which the criteria for required functions performance (output parameters and their allowable ranges) are interrelated, mutually conditioned and dedicated to achieving the specified performance and dependability [44]. Since recently, this approach complies with the logic of the "Space Systems and Complexes" series of standards developed by TsNIIMash in 2019 and 2020.

1. After the introduction of the state standard GOST R 58629, one of the tasks of the failure mode, effects and criticality analysis of space products and processes is aimed at identifying the key design (functional and physical) characteristics of the critical elements and their testability. However, the standard does not establish the method for solving the problem of identification of such key characteristics (although based on the general concept of FMEA, it can be assumed that they are identified, for example, by the method of expert evaluation). Additionally, it is not perfectly clear what should be done if the fulfilment of the required functions cannot be expressed in physical values, but can be characterized with the qualitative features of a product that are described by probabilities (as the degree of confidence that under the specified conditions an event will occur). Nevertheless, in general, the requirements of GOST R

¹ Principle of physicality: The principle, according to which inherent to any system (regardless of its nature) are laws (regularities), perhaps unique, that define the internal causal relations of its existence and operation.

² Physical necessity: The actual causality between a phenomenon and certain natural circumstances that is unambiguously predictable within the knowledge of it (as opposed to randomness).

58629 for the identification of the key characteristics of critical elements comply with the concept of functionality identification in the generalized parametric model of product operation [44].

2. Worst case analysis according to GOST R 58626 allows defining and sets forth a list of formalized analysis procedures that include the quantification of the tolerances of value changes of the output parameters of the object of analysis depending on the possible values of its internal and input parameters. This procedure is the definition of product performance under the worst combinations of the modes and conditions of application in the generalized parametric model of product operation [44]. However, according to GOST R 58626, such analysis is conducted on the basis of the results of FMECA (FMEA) performed in accordance with the requirements of GOST R 58629, i.e., using the method of expert assessment based on experts' opinions for the purpose of subsequent decision-making, which is not a sufficient condition for establishing a complete list of worst cases. Additionally, due to the insufficient maturity of certain terms, for example, "mode" and "emergence" [44, 45], the approach to the worst case analysis according to GOST R 58626 remains uncertain in terms of calculation of the maximum and minimum values of allowable deviations (the worst case) of the output parameters.

3. According to the explanations in the reference annex to GOST 27.002-89, it is not customary to distinguish between the indicator of the probability of no-failure in terms of the strength on the basis of statistical data and the probability of that within the specified period of time the strength values will be within the acceptable limits taking into account the safety factors and strength margins [18]. In fact, this approach to assessing the probability of no-failure corresponds to the definition of dependability taking into account the design margins in such a way as to, with a reliable confidence, guarantee that the values of the examined parameters are within the allowable area [17, 44].

Thus, the key task associated with the identification of possible causes of rare failures consists in conducting a systematic and qualified analysis for identifying the functionality, performance and dependability of a product taking into account critical output parameters and probabilistic indicators that affect the performance of the required functions with the allowable probability of failure. The solution is found using a generalized parametric model of operation and design engineering analysis of dependability.

Ways of achieving systemic analysis

The analysis of the functionality, performance and dependability of products is done based on the information on the modes and conditions of such product's

application, as well as the current state of the design documentation (“as is”) taking into account its requirements for the manufacturing process and technical oversight [5]. The efficiency of such analysis is the highest if it is made on the basis of the intended use, i.e., the key purpose the product is created for that includes (besides the general design goals) all the additional conditions, limitations and requirements that quantify and specify such purpose [46]. After the intended use has been established, a task tree is built for the product’s components that enable the key purpose. Based on each task, the required functions are formulated (defined by the question: “What does an object or its individual elements do?”), each of which is an external manifestation of the product’s properties of a strictly defined physical nature within the given modes and conditions of application. The resulting tree of required functions is the necessary and sufficient condition for substantiating the specified performance and dependability of the product on the basis of the engineering disciplines and methods of ensuring dependability.

After the tree of required functions has been constructed, it becomes possible to identify potential failures in the form of a verbal description of hypothetical events that prevent the performance of the respective functions. Then, the conditions that make failures impossible are defined (failure-free conditions). Such conditions are found using the method of antithesis. The logical design of this method is based on a biased judgement, according to which a failure of any critical element has already “occurred”. If, in the course of design, the required and sufficient measures for eliminating the cause of a possible failure have been taken and documented, that serves as evidence that the above negative judgement is false and, therefore, the condition of reliability has been ensured. The condition of reliability is understood as each of the properties of a particular critical element that makes the corresponding cause of failure impossible [47]. Importantly, under this approach, the properties of critical elements that define their reliability are identified automatically based on strictly engineering techniques (with no regard for the subjective opinion of “experts”).

The list of properties of critical elements itself allows characterizing each critical element quantitatively depending on the selected functional model, i.e., stochastic or physical [17]. Additionally, if the behaviour of a critical element can be characterized in physical values, the description of the properties of a critical element is based on output parameters that best describe the physical nature within the specific system of relations of a specific element in the product and between the entire product and the external environment. This procedure fully complies with the requirements of GOST R 58629. If the model does not allow characterizing the opera-

tion of a critical element through physical values (due to the insufficiency of knowledge regarding the physical nature of failures), the properties of such element are described through an indicator in the form of the probability of failures in the course of performance of the required function based on a “black box” or logical and probabilistic models. If required, the parameters and probabilistic functional indicators of the item can be reduced to a consistent dimensionless form (when the parameters can be represented as the probability of value variation within the allowed range similarly to the explanation given in the Reference Annex to GOST 27.002-89 [18]). That allows estimating the predicted (planned) dependability of the product using the method of dependability calculation based on the probability of performance by components and elements of their required functions [17].

Dependability analysis of highly critical non-recoverable products with short operation life

Based on the above approach, the method of design engineering analysis of dependability (DEAD) [5, 15-17, 36, 37, 44] has been developed, whose application does not cancel any engineering practices, but develops and complements them, enabling the following:

- abandoning the concept of randomness of the causes of failures and establishing their logical and mathematical relationship with the design and engineering factors;
- identifying the relationships between the output operational parameters and the probability of failure;
- identifying the design and manufacturing risks associated with failures that cannot be identified through the conventional methods of analytical and experimental verification;
- timely detecting rare causes of possible failures;
- reducing the number of potential structural failures at early life cycle stages, etc.

Despite the fact that the method of design engineering analysis of dependability (DEAD) implies dependability estimation (calculation), it should be above all considered as a system of design engineering and managerial measures aimed at eliminating (reducing the probability) of failures based on the analysis of the engineering documentation that includes:

- definition of the calculation task (required and sufficient calculations of the performance and dependability parameters according to specified criteria for maximum possible reduction of the probability of unreasonable risks of failures);
- experimental program definition, including experimental identification of the values that cannot be calculated due to the lack of required data and confirmation of the specified performance parameters in the course of

ground experiments, when the number of items submitted for testing is limited due to financial and economic reasons;

- definition of the necessary and sufficient requirements in the design documentation for the manufacture and operation of products;
- development of a check list of output parameters used in the process of quality and dependability verification of products;
- planning of measures to prevent design failures at all life cycle stages;
- iterative calculation of predicted dependability as the result of the required measures to prevent design failures;
- evaluation of design and process engineering solutions for compliance with the specified dependability requirements.

The use of design engineering analysis of dependability (DEAD) creates conditions, in which ensuring dependability is a natural and integral part of a designer's work enabling engineering decision-making in accordance with the specified dependability requirements (rather than in isolation, as it is the case when statistical approaches to dependability are used). However, unlike in the case of failure mode, effects and criticality analysis (FMECA) that is intended for identifying and assessing the criticality of product defects and inadequacies (based on the result of business processes or procedures), DEAD serves to verify designer solutions subject to process constraints (aimed at preventing the causes of possible failures at the physical level before the product has been manufactured).

DEAD has been tested in the design of single-use mechanical space devices and hydraulic assemblies of oil well equipment [5, 15]. Such analysis was in all cases carried out after the experimental activities (and even flight qualification) adopted by the company that developed the mechanisms in accordance with the required regulatory documentation. Nevertheless, the analysis enabled a practically substantiated identification of design and process engineering errors in the design documentation; an evaluation of the effectiveness of the existing computational and experimental optimization of product design; assessment of the adequacy of the established requirements in the design documentation; identification of unacceptable combinations of structural parameters based on the design constraints, actual manufacturing and control conditions; drawing conclusions regarding the products' propensity to failure; predicting the compliance to the specified dependability requirements; providing recommendations regarding design modifications to ensure specified dependability of products. In fact, DEAD allows identifying and eliminating the shortcomings of the conventional methods of design, project engineering and experimental optimization to achieve the specified dependability [5, 15].

Conclusion

For highly critical non-recoverable space entities with short operation life (principally, single-use mechanical devices of spacecraft), the reliability requirements should be considered primarily in terms financial, economic, safety-related and reputational risks associated with the loss of spacecraft. From a design engineer's standpoint, the number of nines after the decimal point (rounded to a smaller number of nines for increased confidence) should be seen as the indicator for the application of the appropriate approaches to ensuring the required reliability at the stage of product design.

In case of two nines after the decimal point it is quite acceptable to use analytical and experimental verification techniques common to the aerospace industry, i.e., dependability calculations using the statistical methods of the modern dependability theory and performance indicators (out of the recommended list according to OST 92-0290), FMEA and Stage-Gate, ground experiments on single workout equivalents for each type of effect [5-8, 27-31, 35, 38-43].

As the required number of nines grows, it is advisable to also use early failure prevention methods, one of which is DEAD that enables designers to adopt substantiated design solutions on the basis of engineering disciplines and design and process engineering methods of ensuring quality and dependability [5, 15-17, 36, 37, 44].

The choice of either of the above dependability strategies is determined solely by the developer's awareness and understanding of potential hazards, which allows managing the risk of potential rare failures or reasonably refusing to do so.

References

1. Fortescue P., Stark J., Swinerd G. Spacecraft systems engineering. NJ: John Wiley & Sons; 2003.
2. Always P. Rockets of the world. Saturn Press; 1999.
3. Sevastyanov N.N. Reliability control in spacecraft with long service life. *Cosmonautics and rocket engineering* 2017;3:133-148. (in Russ.)
4. Gorbenko A.V., Zasukha S.O., Ruban V.I. et al. Safety of rocket-space engineering and reliability of computer systems: 2000-2009 years. *Aerospace technic and technology* 2011;1:9-20. (in Russ.)
5. Pokhabov Yu.P. [Theory and practice of ensuring the dependability of single-use mechanical devices]. Krasnoyarsk: SFU; 2018. (in Russ.)
6. Patraev V.E. [Methods of ensuring and assessing the dependability of long active life spacecraft]. Krasnoyarsk: SibGAU; 2010. (in Russ.)
7. Patraev V.E., Khalimanovich V.I. [Dependability of support spacecraft]. Krasnoyarsk: SibGAU; 2016. (in Russ.)

8. Patraev V.E., Shangina E.A. [Dependability of technical systems of spacecraft]. Krasnoyarsk: SFU; 2019. (in Russ.)
9. Spacecraft [Electronic Resource]. ISS-Reshetnev. [accessed: 15.03.2021]. Available at: iss-reshetnev.ru. (in Russ.)
10. Launches: a database [Electronic Resource]. [Launch vehicles, satellites, airplanes, instruments]. [accessed: 15.03.2021]. Available at: <http://ecorospace.me>. (in Russ.)
11. Saleh J.H., Caster J.-F. Reliability and multi-state failures: a statistical approach. First Edition. NJ: John Wiley & Sons; 2011.
12. Gore B.W. ATR-2009(9369)-1. Critical Clearances in Space Vehicles. The Aerospace Corporation; 2008.
13. Hecht H., Hecht M. Reliability prediction for spacecraft: report prepared for Rome Air Development Center: no. RADC-TR-85-229. Rome Air Development Center; 1985.
14. Tumanov A.V., Zelentsov V.V., Shcheglov G.A. [Fundamentals of spacecraft on-board equipment layout design]. Moscow: Bauman MSTU Publishing; 2010. (in Russ.)
15. [Model and methodology for assessing the impact of scheduled activities for the prevention of structurally defined failures on the dependability of spacecraft in terms of failures of the single-use mechanical devices]. IDN 532-OT-MKB-0031-19. No. GR 1920730200892217000241851. Zheleznogorsk: AO NPO PM MKB; 2019. (in Russ.)
16. Pokhabov Yu.P. Problems of dependability and possible solutions in the context of unique highly vital systems design. *Dependability* 2019;19(1):10-17.
17. Pokhabov Yu.P. Dependability from a designer's standpoint. *Dependability* 2020;4: 13– 20.
18. Annex (informative). GOST 27.002-89. Industrial product dependability. General principles. Terms and definitions. Moscow: Izdatelstvo Standartov; 1990. (in Russ.)
19. Riabinin I.A. [Academy member A.I. Berg and the problems of dependability, survivability and safety]. In: [Academy member Aksel Ivanovich Berg (On the occasion of centenary of the birth)]. Moscow: State Polytechnic Museum; 1993. (in Russ.)
20. Riabinin I.A. [Foundations of the theory and calculation of the dependability of naval power systems]. Leningrad: Sudostroenie; 1971. (in Russ.)
21. Bolotin V.V. [Application of probability theory and dependability theory methods in structural analysis]. Moscow: Stroyizdat; 1971. (in Russ.)
22. Volkov L.I., Shishkevich A.M. [Aircraft dependability]. Moscow: Vyshaya shkola; 1975. (in Russ.)
23. Timashev S.A., Pokhabov Yu.P. [New methods of analysing and evaluating the dependability of aerospace products]. In: [Safety and monitoring of man-made and natural systems: materials and presentations. VI All-Russian Conference (September 18-21, 2018, Krasnoyarsk). Krasnoyarsk: SFU; 2018:254-259. (in Russ.)
24. [Roscosmos to cut launch prices by 30% due to Musk's SpaceX dumping]. RosBusinessConsulting. (accessed 15.03.2021). Available at: https://www.rbc.ru/technology_and_media/10/04/2020/5e90869c9a7947d4640156b7. (in Russ.)
25. Krylov A.M. [Comparative analysis of space activities of Russia, China and India]. MKK. (accessed 15.03.2021). Available at: <http://mosspaceclub.ru/base/base.php>. (in Russ.)
26. Belozertsev A.I., El-Salim S.Z. [Deterministic model of improved dependability of analytical systems]. In: [Dependability and quality: proceedings of the international symposium] 2017;2:396-399. (in Russ.)
27. Kuznetsov A.A. [Structural dependability of ballistic missiles]. Moscow: Mashinostroenie; 1978. (in Russ.)
28. Kuznetsov A.A., Zolotov A.A., Komyagin V.A. et al. [Dependability of mechanical parts of aircraft design]. Moscow: Mashinostroenie; 1979. (in Russ.)
29. Shatrov A.K., Nazarova L.P., Mashukov A.V. [Mechanical devices of spacecraft. Design solutions and dynamic characteristics]. Krasnoyarsk: SibGAU; 2006. (in Russ.)
30. Shatrov A.K., Nazarova L.P., Mashukov A.V. [Introduction to the design of mechanical devices of spacecraft. Design solutions, dynamic characteristics]. Krasnoyarsk: SibGAU; 2009. (in Russ.)
31. Romanov A.V., Testodov N.A. [Introduction to the design of information management and mechanical systems of spacecraft]. Saint Petersburg: Professional; 2015. (in Russ.)
32. Postma R.W. Force and torque margins for complex mechanical systems. In: Proceedings of the 37th Aerospace Mechanisms Symposium, Johnson Space Flight Center, May 19–21; 2004. P. 107-118.
33. Conley P.L., editor. Space vehicle mechanisms – Elements of successful design. NJ: John Wiley & Sons; 1998.
34. Dhillon B.S., Singh C. Engineering reliability. NJ: John Wiley & Sons; 1981.
35. Kolobov A.Yu., Dikoun E.V. Interval estimation of reliability of one-off spacecraft. *Dependability* 2017;4:23-26.
36. Pokhabov Yu.P. Dependability from a designer's standpoint. *Dependability* 2020;2:3-11.
37. Pokhabov Yu.P. Designing complex products with small probability of failure in the context of Industry 4.0. *Ontology of Designing* 2019;9(1):24-35. (in Russ.)
38. [Milestones in space: collection of research papers dedicated to the 50-th anniversary of ISS-Reshetnev]. Krasnoyarsk: IP Sukhovolskaya Yu.P.; 2009. (in Russ.)

39. Crowe D. et al., editors. Design for dependability. NJ: CRC PRESS LLC; 2001.

40. Doronin S.V., Pokhabov Yu.P. [Approaches to the selection of loading cases of load-bearing structures of technical objects]. In: Moskvichiov V.V., editor. [Safety and monitoring of natural and man-made systems: materials and presentations. VII All-Russian Conference (October 5-9, 2020, Kemerovo)]. Novosibirsk: FRC ICT; 2020. P. 43-46. (in Russ.)

41. Pistsova Yu.P., Nikolaeva N.G., Priymak E.V. et al. [Failure mode and effects analysis (FMEA) of design documentation]. [*Bulletin of the Kazan Technological University*] 2004;1:411-415. (in Russ.)

42. Isaev S.V. [We don't need such FMEA! (Difficulties of deployment and infancy mistakes)]. [*Quality management methods*] 2008; 3:30-32. (in Russ.)

43. Stengach M.S., Gorbunov A.A., Kobzev V.N., compilers. [Failure (defect) mode and effects analysis, FMEA]. Samara; 2011.

44. Pokhabov Yu.P. Design for dependability highly responsible systems on the example of a moving rod. *J. Sib. Fed. Univ. Eng. technol.* 2019;12(7). 861-883. (in Russ.)

45. Tarasenko F.P. [Applied systems analysis]. Moscow: KNORUS; 2017. (in Russ.)

46. Baranchukova I.M., Gusev A.S., Kramarenko

Yu.B. et al. [Designing automated machine-building processes]. Moscow: Vysshaya Shkola; 1999. (in Russ.)

47. Gersevanov N.M. [Application of mathematical logic to structural analysis. Volume 1]. Moscow: Stroy-voenmorizdat; 1948. (in Russ.)

About the author

Yuri P. Pokhabov, Candidate of Engineering, Joint Stock Company NPO PM – Maloe Konstruktorskoye Buro (OAO NPO PM MKB), Head of Research and Development Center, Zheleznogorsk, Krasnoyarsk Krai, Russian Federation, e-mail: pokhabov_yury@mail.ru

The author's contribution

The paper considers the matters associated with ensuring the dependability of highly critical non-recoverable entities with short operation life based on one of the methods of early prevention of structurally-defined failures. The paper builds upon the author's ideas set forth in the Dependability Journal nos. 2 and 4, 2020.

Conflict of interests

The author declares the absence of a conflict of interests.

Method of estimating the size of an SPTA with a safety stock

Valery A. Chepurko^{1*}, Alexey N. Chernyaev²

¹RASU, Moscow, Russian Federation, ²MPEI, Moscow, Russian Federation

*VAChepurko@rasu.ru



Valery A. Chepurko



Alexey N.
Chernyaev

Abstract. Aim. To modify the classical method [1, 4] that causes incorrect estimation of the required size of SPTA in cases when the replacement rate of failed parts is comparable to the SPTA replenishment rate. The modification is based on the model of SPTA target level replenishment. The model considers two situations: with and without the capability to correct requests in case of required increase of the size of replenishment. The paper also aims to compare the conventional and adjusted solution and to develop recommendations for the practical application of the method of SPTA target level replenishment. **Methods.** Markovian models [2, 3, 5] are used for describing the system. The flows of events are simple. The final probabilities were obtained using the Kolmogorov equation. The Kolmogorov system of equations has a stationary solution. Classical methods of the probability theory and mathematical theory of dependability [6] were used. **Conclusions.** The paper improves upon the known method of estimating the required size of the SPTA with a safety stock. The paper theoretically substantiates the dependence of the rate of backward transitions on the graph state index. It is shown that in situations when the application is not adjusted, the rates of backward transitions from states in which the SPTA safety stock has been reached and exceeded should gradually increase as the stock continues to decrease. The multiplier will have a power-law dependence on the transition rate index. It was theoretically and experimentally proven that the classical method causes SPTA overestimation. Constraint (3) was theoretically derived, under which the problem is solved sufficiently simply using the classical methods. It was shown that if constraint (3) is not observed, mathematically, the value of the backward transition rate becomes uncertain. In this case, correct problem definition results in graphs with a linearly increasing number of states, thus, by default, the problem falls into the category of labour-intensive. If the limits are not observed, a simplifying assumption is made, under which a stationary solution of the problem has been obtained. It is shown that, under that assumption, the solution of the problem is conservative. It was shown that, if the application is adjusted, the rate of backward transition from the same states should gradually decrease as the stock diminishes. The multiplier will have a hyperbolic dependence on the transition rate index. This dependence results in a conservative solution of the problem of replenishment of SPTA with application adjustment. The paper defines the ratio that regulates the degree of conservatism. It is theoretically and experimentally proven that in such case the classical method causes SPTA underestimation. A stationary solution of the problem of SPTA replenishment with application adjustment has been obtained. In both cases of application adjustment reporting, a criterion has been formulated for SPTA replenishment to a specified level. A comparative analysis of the methods was carried out.

Keywords: Markovian analysis, graph, graph state, probability of transition, failure rate, replenishment rate, SPTA, safety stock, application adjustment.

For citation: Chepurko V.A., Chernyaev A.N. Method of estimating the size of an SPTA with a safety stock. *Dependability* 2021;3:13-19. <https://doi.org/10.21683/1729-2646-2021-21-3-13-19>

Received on: 31.03.2021 / **Upon revision:** 22.07.2021 / **For printing:** 17.09.2021

Introduction

The sufficiency indicator (SI) is the primary quantitative characteristic of SPTA stock. The proposed methods for estimating the size of SPTA usually assume that SI is defined in the performance specifications (PS) for the development of a product or SPTA kit.

An estimate of a stock type is understood as the identification of the required size of SPTA, under which the required SI value is attained. The stock estimation of an SPTA kit consists of the estimates of each individual stock type.

There are several strategies of SPTA replenishment:

- scheduled replenishment (SR). In case of SR, the SPTA replenishment period T_R is defined, after which the stock is replenished to the initial level. If the stock is exhausted before that moment and another failure occurs, the system goes into an inoperable state until the next replenishment;

- scheduled replenishment with emergency deliveries (RED). For the case of RED, the period T_R is also defined. At the same time, upon the SPTA exhaustion and after another failure, an early (emergency) delivery of spare parts (SP) is organized for the purpose of replacing the failed module and replenishing the stock to the initial level;

- continuous replenishment (CR). In case of CR, a stock replenishment request is generated and submitted for execution after each failure and use of each SP. SPTA replenishment may be interpreted as repairs of a failed SP, that after a 100% recovery is added to the SPTA;

- replenishment to the specified level (RL). In case of RL, a stock replenishment application is generated after the stock has reduced to the specified level, including zero, even before the next failure occurs. An application is always generated when the system is operational. That is the difference between RL and RED. This strategy is examined and modified in this paper.

The required size of SPTA is normally estimated using Markovian analysis according to GOST R 51901.15.

The assumptions of the Markovian analysis regarding the probability of transition can be defined as follows:

- state transitions are statistically independent events;
- failure rate λ and recovery rate μ are constant;
- the probabilities of transition from one state to another within a small period of time Δt (Δt is little) are defined by the values $\lambda\Delta t$ and/or $\mu\Delta t$;
- the conditions and mode of operation of all same-type modules are assumed to be identical.

Let us examine the following mathematical model that allows calculating the required SPTA size under condition

that a certain critical level m of safety stock is reached [1, 4]. The replacement process with no repair of failed components corresponds to the random death and birth process that is described by the following state graph (Fig. 1). The system state graph has $k + 2$ states.

Let us introduce the following designations:

n is the number of same-type elements in the system;

k is the planned SPTA size;

m is the level of safety stock, that, when reached, causes the next SPTA replenishment to the planned value k ($1 \leq m \leq k$);

λ is the failure rate of one element;

$\mu = 1/T_D$ is the rate of scheduled SPTA replenishment;

T_D is the average time of SP delivery from the moment the replenishment application is generated.

On the graph, the numbers 0, 1, 2, ..., k show the operable states corresponding to the expended share of the SPTA. State $k + 1$ indicates that the SPTA has been exhausted. The arrows represent transitions from one state to another. Above the arrows are the corresponding transition rates.

Case 1. The application is not adjusted

First, let us assume that the delivery application is not adjusted (see Fig. 1) if further failures occur prior to the delivery of the SP. Let us identify the transition rates for cases when the SP replenishment application has been sent and additional failures have occurred.

Obviously, $\gamma_0 = 1$. The time of transition from state $k - m + 1$ to state 1 will be less than T_D , as the SP replenishment application was previously submitted in state $k - m$. The proportion of the average transition time from state $k - m + 1$ to state 1 relative to the total average time of consecutive transition

from state $k - m$ to state $k - m + 1$, then to state 1 is $\frac{1/\mu}{1/\mu + 1/n\lambda}$

. Thus, the average time of transition from state $k - m + 1$ to state 1 will be equal to $\frac{(1/\mu)^2}{1/\mu + 1/n\lambda}$ and its rate, respectively,

equal to $(1 + \rho)\mu$, where $\rho = \frac{\mu}{n\lambda}$. I.e., the multiplier $\gamma_1 = 1 + \rho$ (see Fig. 1).

Let us summarize. Let us examine a transition from state $k - m + j$ to state j . The transition is possible if the random time η of SPTA replenishment application completion turned out to be longer than the $\xi_1 + \dots + \xi_j = \sigma_j$ total operation time. The average transition time (in terms of mathematical expectation) will be defined by the following integral:

$$E[(\eta - \sigma_j)I\{\eta > \sigma_j\}] = \int_0^{\infty} f_{\sigma_j}(x) \int_x^{\infty} (y - x) f_{\eta}(y) dy dx, \quad (1)$$

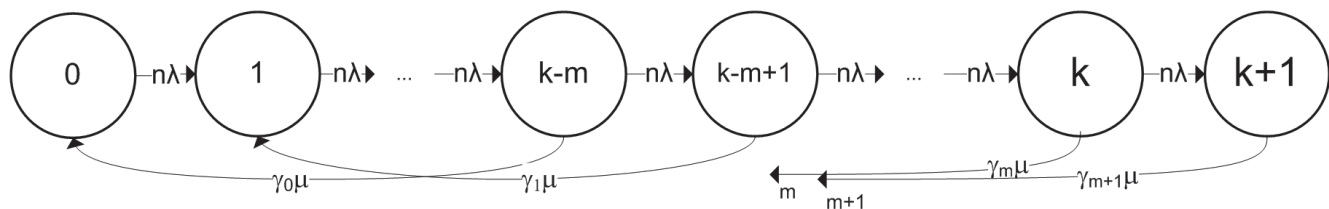


Fig. 1. Transition graph of the death and birth process

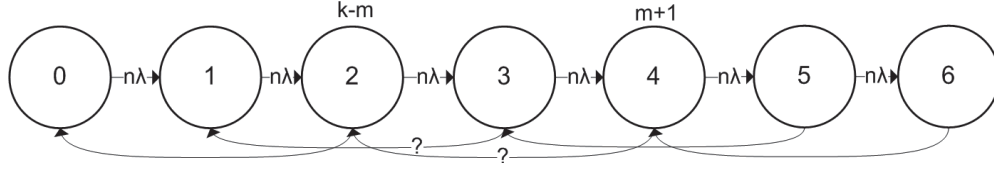


Fig. 2. Transition graph with uncertain transition rate

where $E[\cdot]$ is the expectation operator, $I\{A\} = \begin{cases} 1, & \text{if } A \text{ is true} \\ 0, & \text{if } A \text{ is false} \end{cases}$ is the indicator function, $f_{\sigma_j}(x)$ is the distribution density of the total operation time, $f_{\eta}(y)$ is the distribution density of application execution time. As $\xi_1, \dots, \xi_j$ are independent, identically distributed random values that follow an exponential distribution with the rate of $n\lambda$, then the total operation time will have gamma distribution with the rate of $n\lambda$ and shape variable j . By substituting the exponential probability density function with the rate of j into (1) as $f_{\eta}(y)$, after some simple transformations, the following result is obtained:

$$E\left[(\eta - \sigma_j) I\{\eta > \sigma_j\}\right] = \frac{1}{\mu} \left(\frac{n\lambda}{n\lambda + \mu} \right)^j.$$

Therefore, coefficients γ_j (see Fig. 1) will be equal to:

$$\gamma_j = (1 + \rho)^j, \text{ where } \rho = \frac{\mu}{n\lambda}. \quad (2)$$

The rate of transition from state $k-m+j$ into state j will be $\gamma_j \mu$.

The process of death and birth can be described with a Kolmogorov system of equations. A number of situations is possible involving various numbers of incoming and outgoing transitions.

First, let us introduce a supplementary condition:

$$k \geq 2m + 1. \quad (3)$$

This restriction is to prevent uncertainty in the transition rates (see Fig. 2). Input data: $k=5, m=3$. On the one hand, the rate of transition from state 3 to state 1 should be $\gamma_1 \mu$ if the system has transitioned from state 2 to state 3. On the other hand, if the system has transitioned from state 5 into state 3 and the SPTA was replenished as a result, the rate should be $\gamma_0 \mu$. The similar uncertainty characterizes the rate of transition from state 4 to state 2. It is either $\gamma_2 \mu$ or $\gamma_0 \mu$. The constraint (3) arises from the solution of inequality $k-m \geq m+1$, a condition under which no uncertainty arises (see Fig. 2).

In total, there may be five cases.

1. $k \geq 2m+3$,

$$\begin{cases} P'_0(t) = -n\lambda P_0(t) + \gamma_0 \mu P_{k-m}(t); \\ P'_i(t) = n\lambda P_{i-1}(t) - n\lambda P_i(t) + \gamma_i \mu P_{k-m+i}(t), 1 \leq i \leq m+1; \\ P'_i(t) = n\lambda P_{i-1}(t) - n\lambda P_i(t), m+2 \leq i \leq k-m-1; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_{i-k+m} \mu) P_i(t), k-m \leq i \leq k; \\ P'_{k+1}(t) = n\lambda P_k(t) - \gamma_{m+1} \mu P_{k+1}(t). \end{cases}$$

2. $k=2m+2$,

$$\begin{cases} P'_0(t) = -n\lambda P_0(t) + \gamma_0 \mu P_{k-m}(t); \\ P'_i(t) = n\lambda P_{i-1}(t) - n\lambda P_i(t) + \gamma_i \mu P_{k-m+i}(t), 1 \leq i \leq m+1; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_{i-k+m} \mu) P_i(t), m+2 \leq i \leq k; \\ P'_{k+1}(t) = n\lambda P_k(t) - \gamma_{m+1} \mu P_{k+1}(t). \end{cases}$$

3. $k=2m+1$,

$$\begin{cases} P'_0(t) = -n\lambda P_0(t) + \gamma_0 \mu P_{k-m}(t); \\ P'_i(t) = n\lambda P_{i-1}(t) - n\lambda P_i(t) + \gamma_i \mu P_{k-m+i}(t), 1 \leq i \leq m; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_0 \mu) P_i(t) + \gamma_i \mu P_{k-m+i}(t), i = m+1; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_{i-k+m} \mu) P_i(t), m+2 \leq i \leq k; \\ P'_{k+1}(t) = n\lambda P_k(t) - \gamma_{m+1} \mu P_{k+1}(t). \end{cases}$$

4. $m+2 \leq k \leq 2m$,

$$\begin{cases} P'_0(t) = -n\lambda P_0(t) + \gamma_0 \mu P_{k-m}(t); \\ P'_i(t) = n\lambda P_{i-1}(t) - n\lambda P_i(t) + \gamma_0 \mu P_{k-m+i}(t), 1 \leq i \leq k-m-1; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_0 \mu) P_i(t) + \\ + \gamma_{(k-2m-1+i)} \mu P_{k-m+i}(t), k-m \leq i \leq m+1; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_{i-m-1} \mu) P_i(t), m+2 \leq i \leq k; \\ P'_{k+1}(t) = n\lambda P_k(t) - \gamma_{k-m} \mu P_{k+1}(t). \end{cases}$$

5. $k=m+1$,

$$\begin{cases} P'_0(t) = -n\lambda P_0(t) + \gamma_0 \mu P_{k-m}(t); \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_0 \mu) P_i(t) + \gamma_0 \mu P_{k-m+i}(t), 1 \leq i \leq m; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_0 \mu) P_i(t) + \gamma_1 \mu P_{k-m+i}(t), i = m+1; \\ P'_{k+1}(t) = n\lambda P_k(t) - \gamma_1 \mu P_{k+1}(t). \end{cases}$$

Let us search for a stationary solution. For that purpose, let us assume the derivatives on the left sides of the equations equal 0 and apply the normalization condition: $\sum_{i=0}^{k+1} P_i = 1$.

1. $k \geq 2m+3$,

$$\begin{cases} P_0 = \rho P_{k-m}; \\ P_{i-1} = P_i - \gamma_i \rho P_{k-m+i}, 1 \leq i \leq m+1; \\ P_{i-1} = P_i, m+2 \leq i \leq k-m-1; \\ P_{i-1} = (1 + \gamma_{i-k+m} \rho) P_i, k-m \leq i \leq k; \\ P_k = \gamma_{m+1} \rho P_{k+1}. \end{cases}$$

2. $k=2m+2$,

$$\begin{cases} P_0 = \rho P_{k-m}; \\ P_{i-1} = P_i - \gamma_i \rho P_{k-m+i}, 1 \leq i \leq m+1; \\ P_{i-1} = (1 + \gamma_{i-k+m} \rho) P_i, m+2 \leq i \leq k; \\ P_k = \gamma_{m+1} \rho P_{k+1}. \end{cases}$$

3. $k=2m+1$,

$$\begin{cases} P_0 = \rho P_{k-m}; \\ P_{i-1} = P_i - \gamma_i \rho P_{k-m+i}, 1 \leq i \leq m; \\ P_{i-1} = (1 + \rho) P_i - \gamma_i \rho P_{k-m+i}, i = m+1; \\ P_{i-1} = (1 + \gamma_{i-k+m} \rho) P_i, m+2 \leq i \leq k; \\ P_k = \gamma_{m+1} \rho P_{k+1}. \end{cases}$$

4. $m+2 \leq k \leq 2m$,

$$\begin{cases} P_0 = \rho P_{k-m}; \\ P_{i-1} = P_i - \rho P_{k-m+i}, 1 \leq i \leq k-m-1; \\ P_{i-1} = (1 + \rho) P_i - \gamma_{(k-2m-1+i) \vee 0} \rho P_{k-m+i}, k-m \leq i \leq m+1; \\ P_{i-1} = (1 + \gamma_{i-m-1} \rho) P_i, m+2 \leq i \leq k; \\ P_k = \gamma_{k-m} \rho P_{k+1}. \end{cases}$$

5. $k=m+1$,

$$\begin{cases} P_0 = \rho P_1; \\ P_{i-1} = (1 + \rho) P_i - \rho P_{i+1}, 1 \leq i \leq m; \\ P_m = (1 + \rho) P_{m+1} - \gamma_1 \rho P_{m+2}; \\ P_{m+1} = \gamma_1 \rho P_{m+2}. \end{cases}$$

Next, for the first case, let us express the obtained recursions and the final solution. For convenience, let us denote

$$\rho = \frac{\theta}{n} = \frac{\mu}{n\lambda}.$$

$k \geq 2m+3$. The solution is found “from top to bottom”.

$$P_k = \gamma_{m+1} \rho P_{k+1};$$

$$P_{i-1} = \gamma_{m+1} \rho P_{k+1} \prod_{j=m-k+i}^m (1 + \gamma_j \rho), k-m \leq i \leq k;$$

$$P_{i-1} = P_i, m+2 \leq i \leq k-m-1.$$

Therefore,

$$P_{m+1} = P_{m+2} = \dots = P_{k-m-1} = \gamma_{m+1} \rho P_{k+1} \prod_{j=0}^m (1 + \gamma_j \rho) = A_m P_{k+1},$$

where $A_m = \gamma_{m+1} \rho \prod_{j=0}^m (1 + \gamma_j \rho)$.

General formula

$$P_i = \left(A_m - \gamma_{m+1} \rho \prod_{j=i+1}^m (1 + \gamma_j \rho) \right) P_{k+1}, 0 \leq i \leq m-1.$$

Additionally,

$$P_m = (A_m - \gamma_{m+1} \rho) P_{k+1}.$$

After applying the normalization condition, we will obtain the following result

$$(k-m)A_m + 1 = \frac{1}{P_{k+1}}.$$

The second and third situations cause similar conclusions.

Thus, under restriction (3), the probability of SPTA failure will be defined by the expression:

$$P_{k+1} = \frac{1}{1 + (k-m)A_m}, \text{ where } A_m = \gamma_{m+1} \rho \prod_{j=0}^m (1 + \gamma_j \rho). \quad (4)$$

Now, let us drop the condition (3), i.e., let us examine cases 4 and 5 and implement the accurate model of SP supply using a Markovian graph (Fig. 3). The graph, due to the awkwardness that increases as k grows, is represented for the special case $k=3, m=2$.

First, let us consider case 4. In order to avoid uncertainty in the rate of recurrent transitions, additional states have been introduced into SPTA replenishment: (2.0), (2.1), (3.0), (3.1), (3.2), The first number is the number of SP taken from the SPTA, the second one is the number of failures that occurred since the submission of the replen-

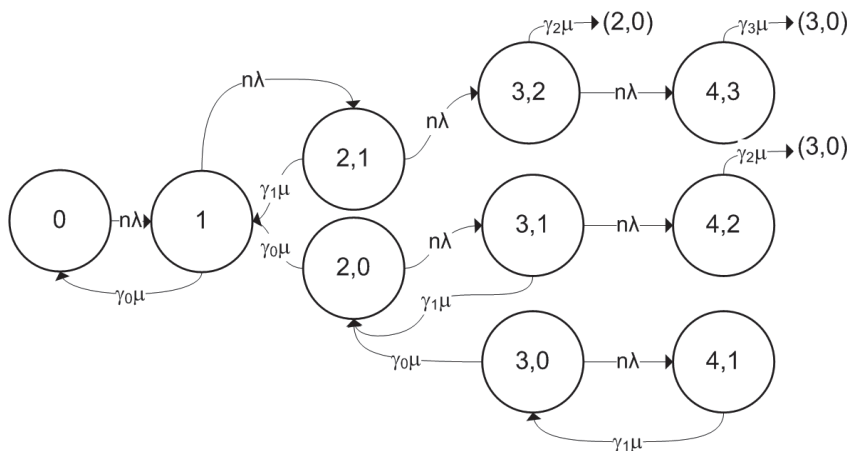


Fig. 3. System transition graph in case of $k \leq 2m$

ishment application. For example, (2.0) means that there have been two failures (and 2 replacements, respectively) and there are no unfilled applications, and it is time to send a new SPTA $k-m=1$ replenishment application. (3.1) means that there have been three failures (3 replacements, respectively), there is an unsatisfied application submitted in a situation where two SPs from the SPTA have been used and it is time to send a new SP $k-m=1$ replenishment application. (3.2) means that there have been three failures (3 replacements, respectively), there is an unsatisfied application submitted in a situation where another SP from the SPTA has been used and it is time to send a new SP $k-m=1$ replenishment application.

Finding an accurate analytical solution in the general case is a time-consuming task.

Therefore, we will solve its simplified version, assuming that in case of uncertainty the multiplier takes the minimal possible values, i.e., γ_0 .

$$\gamma_i = \min(\gamma_0, \gamma_{i,j}) = \gamma_0 = 1. \quad (5)$$

That assumption leads to a conservative estimate of the required SPTA size, since in the simplified version the rate of recurrent transitions into “improved” states decreases, the response to applications slows down. As before, we will find stationary solutions using the same notations

$$\begin{cases} P_0 = \rho P_{k-m}; \\ P_{i-1} = P_i - \rho P_{k-m+i}, 1 \leq i \leq k-m-1; \\ P_{i-1} = (1+\rho)P_i - \gamma_{(k-2m-1+i)\vee 0} \rho P_{k-m+i}, k-m \leq i \leq m+1; \\ P_{i-1} = (1+\gamma_{i-m-1}\rho)P_i, m+2 \leq i \leq k; \\ P_k = \gamma_{k-m} \rho P_{k+1}. \end{cases}$$

The solution of system (6) in an explicit form is quite cumbersome. In order to solve the problem numerically, it is suggested using a replacement-based algorithm

$$P_i = a_i P_{k+1}, i = 0, 1, \dots, k+1. \quad (7)$$

Coefficients a_i will be defined by a “backward” recursion:

$$\begin{cases} a_{k+1} = 1; \\ a_k = \gamma_{k-m}\rho; \\ a_{i-1} = (1+\gamma_{i-m-1}\rho)a_i, m+2 \leq i \leq k; \\ a_{i-1} = (1+\rho)a_i - \gamma_{(k-2m-1+i)\vee 0} \rho a_{k-m+i}, k-m \leq i \leq m+1; \\ a_{i-1} = a_i - \rho a_{k-m+i}, 1 \leq i \leq k-m-1. \end{cases}$$

In (6) and (8), for the purpose of reduction, symbol $\vee$ is used that means $a \vee b = \max(a, b)$. The probability of SPTA shortage will equal

$$P_{k+1} = \left(\sum_{i=0}^{k+1} a_i \right)^{-1}. \quad (9)$$

Let us consider the fifth case: $k=m+1$.

From the first two equations we obtain a recurrence equation:

$$P_{i-1} = \rho P_i, i = 1, \dots, m+1.$$

By using the latter equation, we deduce:

$$P_i = \gamma_1 \rho^{m+2-i} P_{m+2}, i = 0, \dots, m+1.$$

Out of the normalization condition, if $\rho \neq 1$, we obtain:

$$P_{k+1} = \frac{1-\rho}{1-\rho+\gamma_1\rho-\gamma_1\rho^{m+3}} = \frac{1-\rho}{1+\rho^2-(1+\rho)\rho^{m+3}}. \quad (10)$$

If $\rho = 1$, we obtain:

$$P_{k+1} = \frac{1}{1+2(m+2)} = \frac{1}{2m+5}. \quad (11)$$

Case 2. The application is adjusted

Now, let us assume that the delivery application is adjusted in a situation when further failures have occurred before the delivery of an SP kit (Fig. 4).

In this case, the transition rate multipliers will also be non-constant, as additional time will be required for the replenishment per the previous application. Let us suppose that it takes an average time of τ to consolidate an SPTA kit in a warehouse, while the delivery to the destination takes π . Then, the transition from state $k-m$ into state 0 will on average take $m\tau + \pi$. The transition from state $k-m+1$ to state 1 will on average take $(m+1)\tau + \pi$. In general, the transition from state $k-m+j$ into state j will on average take $(m+j)\tau + \pi$. Additionally, $j = 0, 1, \dots, m+1$. As before, $\gamma_0 = 1$. Then,

$$\gamma_j = \frac{m\tau + \pi}{(m+j)\tau + \pi}. \quad (12)$$

If the times τ and π are unknown, we can use the estimate

$$\frac{m\tau + \pi}{(m+j)\tau + \pi} \geq \frac{m}{m+j}.$$

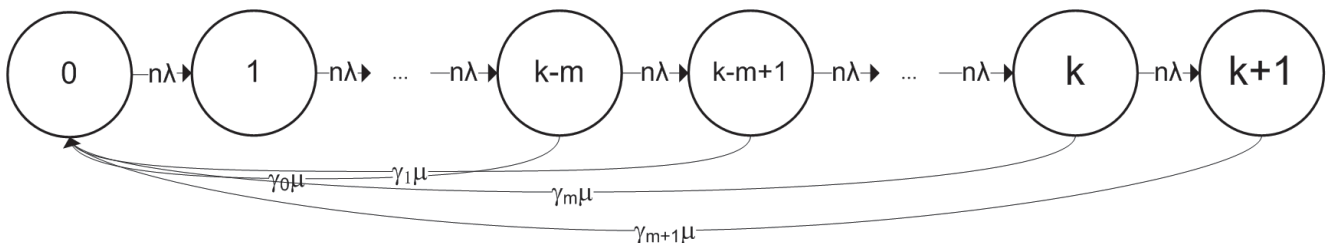


Fig. 4. System transition graph with adjustment of SPTA application

If we assume that

$$\gamma_j = \frac{m}{m+j}, \quad (13)$$

then the recurrent transition rates will be underestimated and the required SPTA size will be overestimated, i.e., conservative. The degree of conservatism will be defined by formula π/τ . If $(\pi/\tau) > 1$, the delivery time is longer than the time of SPTA kit consolidation, the result will be strongly conservative, and an adequate result should be obtained using the classical method of computation. If $(\pi/\tau) < 1$, the delivery time is shorter than the time of single SPTA kit consolidation, the result will be mildly conservative, the error caused by the application of (13) instead of (12) will be small. In any case, if π/τ is known, (12) is preferable.

The following system of differential equations corresponds to the graph shown in Fig. 4.

$$\begin{cases} P'_0(t) = -n\lambda P_0(t) + \mu \sum_{i=0}^{m+1} \gamma_i P_{k-m+i}(t); \\ P'_i(t) = n\lambda P_{i-1}(t) - n\lambda P_i(t), 1 \leq i \leq k-m-1; \\ P'_i(t) = n\lambda P_{i-1}(t) - (n\lambda + \gamma_{i-k+m}\mu)P_i(t), k-m \leq i \leq k; \\ P'_{k+1}(t) = n\lambda P_k(t) - \gamma_{m+1}\mu P_{k+1}(t). \end{cases} \quad (14)$$

As we did before, let us seek a stationary solution “from top to bottom”.

$$P_k = \gamma_{m+1}\rho P_{k+1};$$

$$P_{i-1} = \gamma_{m+1}\rho P_{k+1} \prod_{j=m-k+i}^m (1 + \gamma_j\rho), k-m \leq i \leq k;$$

$$P_{i-1} = P_i, 1 \leq i \leq k-m-1.$$

Therefore,

$$P_1 = P_2 = \dots = P_{k-m-1} = \gamma_{m+1}\rho P_{k+1} \prod_{j=0}^m (1 + \gamma_j\rho) = A_m P_{k+1}.$$

The probability of the original state is defined by the sum

$$P_0 = \rho \sum_{i=0}^{m+1} \gamma_i P_{k-m+i} = \gamma_{m+1}\rho P_{k+1} \left[1 + \gamma_m\rho + \sum_{i=0}^{m-1} \gamma_i\rho \prod_{j=i+1}^m (1 + \gamma_j\rho) \right].$$

After applying the normalization condition, we will obtain the following result

$$\gamma_{m+1}\rho \sum_{i=1}^m \prod_{j=i}^m (1 + \gamma_j\rho) + (k-m)A_m + \gamma_{m+1}\rho + 1 = \frac{1}{P_{k+1}},$$

from which the probability of SPTA shortage will be defined by the formula:

$$P_{k+1} = \frac{1}{(k-m)A_m + \gamma_{m+1}\rho + 1 + B_m}, \quad (15)$$

where $A_m = \gamma_{m+1}\rho \prod_{j=0}^m (1 + \gamma_j\rho)$, $B_m = \gamma_{m+1}\rho \sum_{i=1}^m \prod_{j=i}^m (1 + \gamma_j\rho)$.

As before, we select k so that probability (15) is less than a certain number ε . Finally, we obtain the following criteria for the assessment of the required stock:

$$k \geq \frac{\frac{1-\varepsilon}{\varepsilon} - \gamma_{m+1}\rho - B_m}{A_m} + m, \quad (16)$$

where $A_m = \gamma_{m+1}\rho \prod_{j=0}^m (1 + \gamma_j\rho)$.

Comparative analysis of the obtained results

Let us compare the estimates of the required SPTA size obtained using the classical method with constant rates of recurrent states and the new method taking into account the adjustments γ_i , $i = 1, 2, \dots, m+1$.

Table 1. Estimation of the required size of SPTA. $\rho = 1$. The application is not adjusted

Method	ε		
	0.1	0.05	0.01
Classical	4	6	26
New	3	3	6

Table 2. Estimation of the required size of SPTA. $\rho = 1$. The request is adjusted

Method	ε		
	0.1	0.05	0.01
Classical	3	5	25
New	10	20	100

Table 3. Estimation of the required size of SPTA $\rho = 2$. The application is not adjusted

Method	ε		
	0.1	0.05	0.01
Classical	2	3	7
New	2	2	3

Table 4. Estimation of the required size of SPTA $\rho = 2$. The request is adjusted

Method	ε		
	0.1	0.05	0.01
Classical	2	2	7
New	3	6	26

Table 5. Estimation of the required size of SPTA $\rho = 5$. The application is not adjusted

Method	ε		
	0.1	0.05	0.01
Classical	2	2	2
New	2	2	2

Table 6. Estimation of the required size of SPTA $\rho = 5$. The request is adjusted

Method	ε		
	0.1	0.05	0.01
Classical	1	1	2
New	2	2	4

Tables 1 to 6 show the calculated required SPTA size under various values of $\rho = \frac{\mu}{n\lambda}$ and probabilities of SPTA shortage ε . $m=1$ is taken as the safety stock.

Having analysed the calculation results, we can note that, first, previous conclusions on the overestimation or underestimation of the size of SPTA proved to be correct. Second, the amendments proposed in the paper should be taken into account if $\rho \leq 5$, i.e., when the value of SPTA replenishment rate μ is comparable to the total failure rate $n\lambda$. In this situation, a tangible economic effect can actually be obtained if the application is not adjusted. In the same case, if the application is adjusted, a conservative, yet correct, estimate of the required size of SPTA can be obtained.

On the other hand, it should be noted that the paper examined and investigated the solution of a stationary problem, that will probably significantly differ from the solution of the original non-stationary problem in cases of relatively short calculation time. The authors will dedicate further research activities to the computational solution of a non-stationary problem.

Conclusion

The paper improved the method of estimating the required size of an SPTA with a safety stock. The application results of the adjusted and classical estimates were analysed. It was revealed that the most pronounced effect from the application of the adjusted estimate will be observed if the failure rates are comparable to the SPTA replenishment rates.

References

1. GOST 27.507-2015 Reliability in technique. Spare parts, tools and accessories. Evaluation and calculation of reserves. Moscow: Standartinform; 2017. (in Russ.)

2. GOST R 51901.15-2005 (IEC 61165:1995) Risk management. Application of Markov techniques. Moscow: Standartinform; 2005. (in Russ.)

3. Cherkesov G.N. [Dependability of hardware and software systems: a study guide]. Saint Petersburg: Piter; 2005. (in Russ.)

4. Cherkesov G.N. [System dependability assessment with STPA taken into consideration: a study guide for undergraduate students]. Saint Petersburg: BHV-Peterburg; 2012. (in Russ.)

5. Viktorova V.S., Stepaniants A.S. [Models and methods of technical system dependability calculation]. LENAND (URSS); 2016. (in Russ.)

6. Maksimov Yu.D. [Probability-related branches of mathematics: a textbook for bachelors of engineering]. Saint Petersburg: Ivan Fedorov; 2001. (in Russ.)

About the authors

Valery A. Chepurko, Candidate of Physics and Mathematics, Associate Professor, Chief Specialist of the Division for Computational Substantiation of Design Solutions, RASU, Moscow, Russian Federation, e-mail: VAChepurko@rasu.ru

Alexey N. Chernyaev, Candidate of Engineering, Associate Professor, Head of the Department of Automated Thermal Management Systems, MPEI, Moscow, Russian Federation, e-mail: ChernyaevAN@mpei.ru, phone: +7 (495) 362 77 20, +7 (495) 362 70 29.

The authors' contribution

The authors have modified and expanded the method of evaluating SPTA with safety stock. A comparative analysis of the classical and modified methods was carried out.

Valery A. Chepurko developed a method of assessing the size of SPTA in situations when the SPTA replenishment request is not adjusted.

Alexey N. Chernyaev developed a method of assessing the size of SPTA in situations when the SPTA replenishment request is adjusted and conducted a comparative test.

Conflict of interests

The authors declare the absence of a conflict of interests.

Estimation of the failure flow of a set of passenger car doors

Maria V. Belousova¹, Vitaly V. Bulatov^{2*}, Nikolay V. Smirnov¹

¹ Saint Petersburg State University, Russian Federation, Saint Petersburg

² Saint Petersburg State University of Aerospace Instrumentation, Saint Petersburg, Russia

*bulatov-vitaly@yandex.ru



Maria V. Belousova



Vitaly V. Bulatov



Nikolay V. Smirnov

Abstract. An estimation of the failure flows is a prerequisite for the operation of industrial products. It is based on statistical data about failures that occur within technical items in the process of their operation. In the technical product documentation, this indicator shall be featured in the “Dependability parameter estimation” section. The dependability analysis of rolling stock is still affected by the difficulty of defining the methodology for evaluating this parameter at various system levels. For the purpose of analysing a multicomponent system, a reliability block diagram should be developed, and the possible replacement (redundant) elements should be taken into consideration. Multicomponent systems are often represented through various block diagrams, where, among others, the “m-out-of-n” structure may be used referring to a system with a parallel arrangement of elements that is operable when at least m elements operate. An example of such system is a set of passenger car doors. The manufacturers and customers may have different approaches to calculating technical system dependability. First, the required dependability indicator for the entire train is defined that, in turn, defines the dependability requirements for a car. At the same time, the dependability indicator for a car is determined by the respective values of its components (subsystems, units and parts). However, the nature of the relationship between a car and its components is not always taken into account. At the same time, car manufacturers can and should define in the regulatory documentation (and later supervise in operation) the dependability indicators for a set of doors (components of a car in our case) as a single system. However, the failure criteria of a set of doors are not always defined. This paper examines the method of calculating the failure flow for a set of passenger car doors based on operational data and the failure flow of a single door. **Aim.** To propose a method for calculating the failure flow of a set of 6 car doors by analysing the possible reliability block diagrams with subsequent transition to transition and state graphs. **Conclusions.** A number of block diagrams were developed for the purpose of dependability calculation of sets of passenger car doors based on the system failure criterion. The failure flow of a set of car doors was calculated according to the developed block diagrams. It is concluded that the Markovian method of calculating the failure flow is of higher priority than the logic-and-probability approach, since it takes into account the recovery factor. A Markovian method was proposed for calculating the failure flow and recovery time of a set of car doors for the “3-out-of-4” reliability block diagram.

Keywords: dependability, failure flow, reliability block diagram, set of doors, Markovian processes.

For citation: Belousova M.V., Bulatov V.V., Smirnov N.V. Estimation of the failure flow of a set of passenger car doors. *Dependability* 2021;3: 20-26. <https://doi.org/10.21683/1729-2646-2021-21-3-20-26>

Received on: 20.06.2021 / **Upon revision:** 23.07.2021 / **For printing:** 17.09.2021.

Introduction

One of today's pressing problems consists in ensuring high dependability of railway transportation. The safety of passenger transportation is directly associated with the estimation of the dependability indicators of rolling stock components. The suppliers of passenger car components are to estimate the dependability indicators over the period of time agreed with the customer (car-building plant). The manufacturers, in turn, report to transportation companies with estimated dependability indicators of the products, i.e., the cars. Hence the difference in the approaches to the quantitative estimation:

1) the car manufacturers calculate the dependability indicators of a car, while the suppliers calculate those of the components and elements;

2) the method of estimating the dependability indicators for cars differs from that for rolling stock components due to the different levels of operation and design. That is expressed in the fact that the estimation of the mean operation time of a car could be done using the exponential distribution law, while that of steps, for instance, depending on the version, using both exponential and normal, as well as the Weibull distribution [1]. The exponential distribution is typical for complex items consisting of many elements with different operation time distributions [2]. Since a set of car doors consists of mechanical (manual) and automatic (electromechanical) side and gangway doors, the exponential distribution law should be used when estimating the dependability indicators. The assumption of the Markovian nature

of transitions within a complex system is due to the fact that if each of a systems' elements has an approximately exponential law of fault-free operation distribution, then the behaviour of the entire system can be described with a Markovian process [3]. An automatic door kit includes a programmable control unit designed for opening/closing doors, as well for processing information generated by inductive displacement sensors and obstruction-in-the-opening sensors. Additionally, the failure flow of doors is assumed to be constant in the course of normal operation in accordance with the results of the criterion's application according to GOST R IEC 60605-6-2007. Consequently, it should be assumed that a door's time to failure is exponentially distributed, and it was decided to use Markovian analysis for calculating the failure flow of a set of car doors.

The passenger car manufacturers take into consideration the dependability indicators for entire sets of car doors (4 to 6, depending on the model) according to the following rule: the door's failure flow parameter is multiplied by the number of such doors in the car, which, in terms of structural dependability, indicates an elementary serial structure, where the failure of one door in most cases equals the failure of the car, which is a questionable assumption.

For the purpose of correctly defining the operating procedure of doors as part of a car subject to the hierarchy of connections, the structural approach was considered. A dependability block diagram is a graphical representation of the operational state of a system. It shows the logical connections between the operating components required

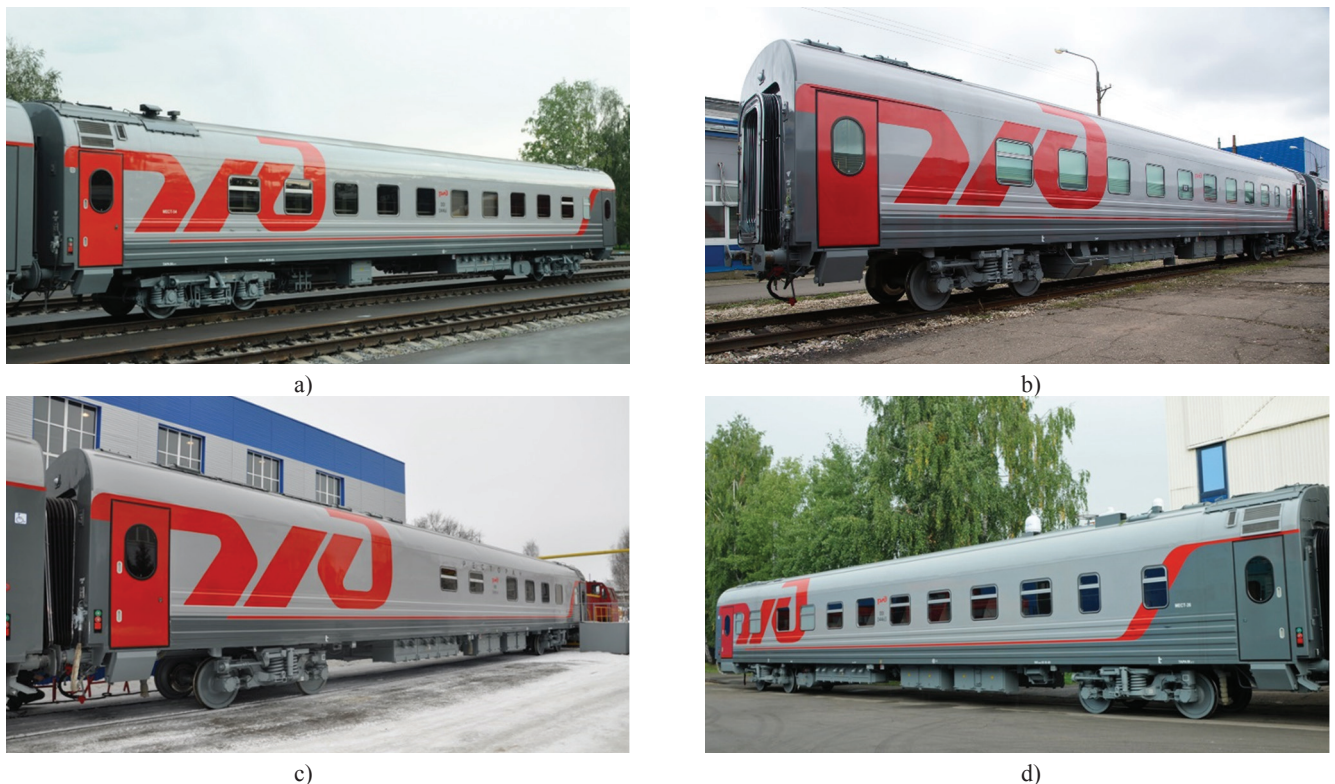


Fig. 1. The models of passenger cars examined in the context of reliability block diagrams are as follows: a) model 61-4447 (couchette car); b) model 61-4462 (compartment car); c) model 61-4460 (dining car); d) model 61-4445 (staff car)

for the system to perform as expected [4]. However, the reliability block diagram-based simulation methods are designed for non-recoverable systems, where the failure order is irrelevant. In the case of recoverable facilities and systems, where the failure order is to be taken into consideration, such methods as the Markovian analysis [5-8] are more appropriate. In the course of the study, the operation of an entire system was represented using the structural approach that enabled – taking into account the restoration of the system – the transition to the state and transition graph [4].

Object of research

Exterior doors are designed for all types of passenger cars with the design speed of up to 200 km/h (Fig. 1). The doors ensure comfort and safety: gangway doors allow moving from car to car; side doors allow entering and exiting a car, protect against sudden changes of pressure and temperature, prevent dust and moisture from entering a car's vestibule, provide noise and heat insulation of a car vestibule under all modes of train operation.

Fig. 2 shows the location of different types of doors of a complete set.

Reliability block diagram of a set of car doors

A reliability block diagram of a set of car doors can be built in a number of ways. In terms of model analysis, it is assumed that in the examined cases the system is under complete control, while the switches are perfectly dependable. The simplest representation of a set of car doors is a non-redundant system (Fig. 3), where the failure of any element causes the failure of the entire system. Then, the system's probability of no-failure is calculated using a well-known formula:

$$P_{cl} = \prod_{i=1}^N P_i, \quad (1)$$

where P_i is the probability of no-failure of the element; N is the number of elements in the system.

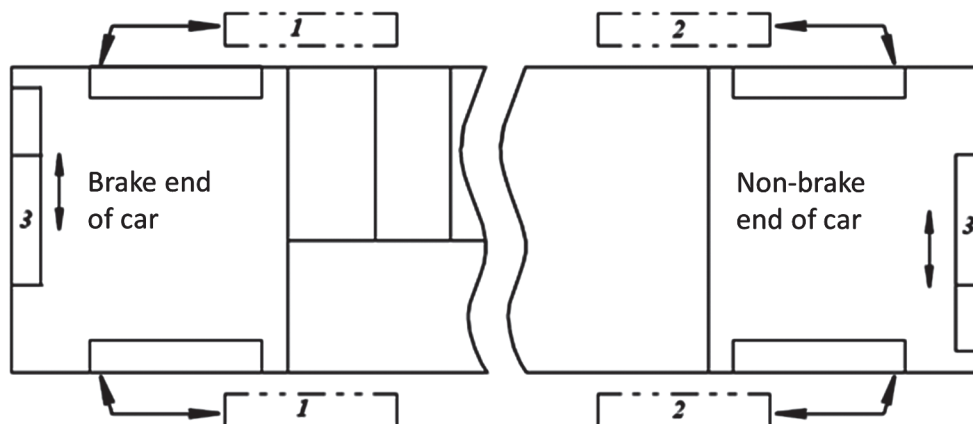


Fig. 2. Passenger car doors: 1 – single side door with electromechanical operation; 2 – single side door with mechanical operation; 3 – single gangway door with electromechanical operation.



Fig. 3. Linear reliability block diagram of a set of passenger car doors, where E_i is a door with electromechanical operation; B_i is a gangway door

The second possible representation of the reliability block diagram of a set of car doors is segregated constant redundancy with integer multiplicity. The various methods of redundancy and their respective dependability benefits are discussed in detail in [9]. The methods of assessing the impact of maintenance on the efficiency of redundancy are described in [10]. In general, the probability of no-failure of this diagram is calculated using the formula:

$$P_p = 1 - \prod_{i=1}^N (1 - P_i). \quad (2)$$

For a set of car doors, let us consider composite diagrams:

- manual door backed-up by another manual door (Fig. 4a);
- automatic side doors backed-up by manual doors (Fig. 4b).

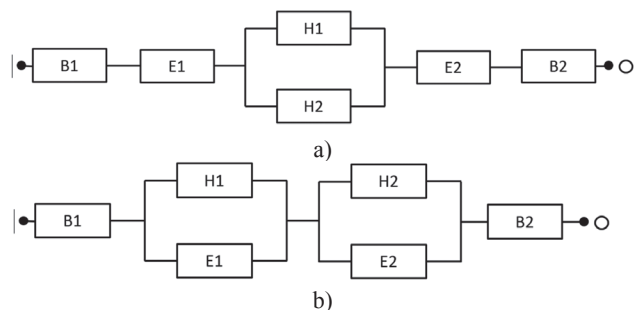


Fig. 4. Redundant reliability block diagrams of a set of passenger car doors, where E_i is a side door with electromechanical operation, B_i is a gangway door, H_i is a door with mechanical operation

Thus, we obtain two variants of a mixed system with element redundancy of individual units.

The probability of no-failure with a back-up manual door (Fig. 4a) is calculated using formulas (1) and (2):

$$P_{c2} = P_{B1} \cdot P_{E1} \cdot P_{B2} \cdot P_{E2} \cdot [1 - (1 - P_{H1}) \cdot (1 - P_{H2})].$$

The probability of no-failure with manual doors backing-up electromechanical side doors (Fig. 4b) is calculated as follows:

$$P_{c3} = P_{B1} \cdot P_{B2} \cdot [1 - (1 - P_{E1}) \cdot (1 - P_{H1})] \cdot [1 - (1 - P_{E2}) \cdot (1 - P_{H2})].$$

Another possible approach to structural evaluation involves representing a set of car doors as a “ m -out-of- n ” structure [4]. A system of this type can be considered a variant of a system with a parallel arrangement of elements that is operable when at least m elements out of n ($m < n$) are operable.

Let us describe three variants:

- a “3-out-of-4” system that is considered operable when doors H1, E1 and E2 are operable (Fig 5a);

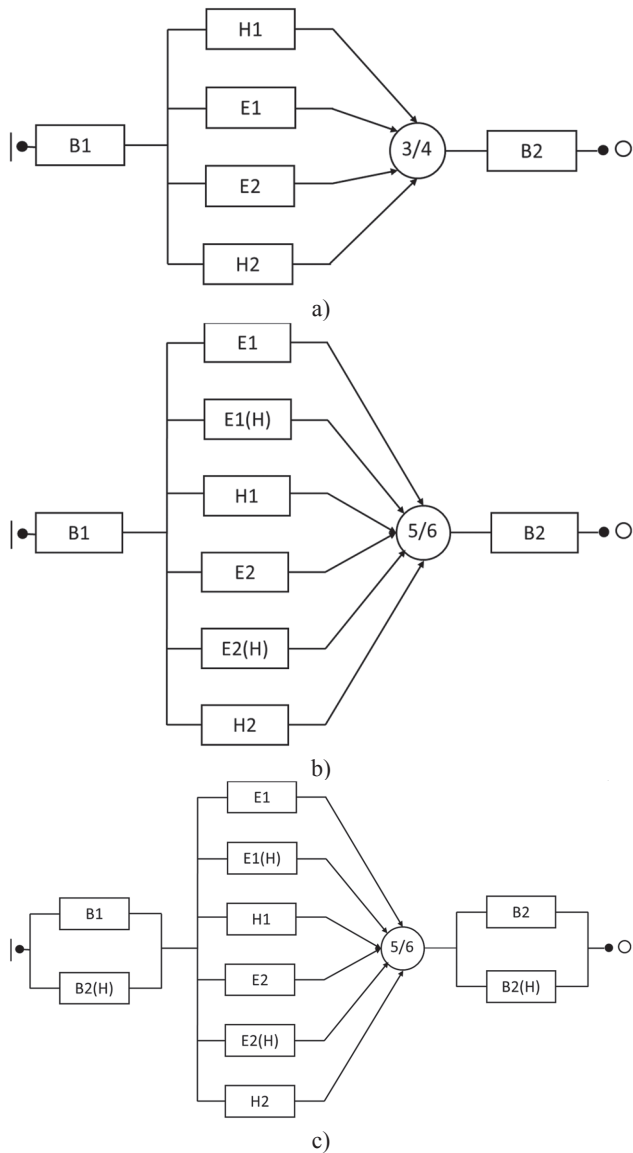


Fig. 5. “ M -out-of- n ” block diagrams for evaluating the dependability of a set of passenger car doors, where Ei is a door with electromechanical operation, Ei(H) is a door with electromechanical operation in manual mode, Bi is a gangway door, Bi(H) is a gangway door with electromechanical operation in manual mode, Hi is a door with mechanical operation

- a “5-out-of-6” system that implies the option of manual operation for electromechanical doors Ei(H) (Fig. 5b);

- a “5-out-of-6” system that takes into account the possible manual operation of electromechanical side and gangway doors, Ei(H) and Bi(H) (Fig. 5b).

Both of the above structures are mixed. Sequentially connected or element-redundant gangway doors are added to the “ m -out-of- n ” structure.

Then, provided that all elements of the “ m -out-of- n ” are equally dependable, the probability of no-failure of the structure shown in Fig. 5a will be as follows:

$$P_{c4} = P_{B1} \cdot P_{B2} \cdot (4 \cdot P_m^3 - 3 \cdot P_m^4),$$

where P_m is an element of the “ m -out-of- n ” structure.

The probability of no-failure for the structure shown in Fig. 5b under the same conditions is as follows:

$$P_{c5} = P_{B1} \cdot P_{B2} \cdot (6 \cdot P_m^5 - 5 \cdot P_m^6).$$

Finally, the probability of no-failure of the structure shown in Fig. 5c that implies both automatic and manual operation of a gangway door will be as follows:-

$$P_{c6} = [1 - (1 - P_{B1}) \cdot (1 - P_{B1(H)})] \cdot [1 - (1 - P_{B2}) \cdot (1 - P_{B2(H)})] \cdot (6 \cdot P_m^5 - 5 \cdot P_m^6).$$

Failure flow calculation for the block diagrams of car door sets

Let us estimate the failure flow for the structures shown in Fig. 4b, 5a-c. The mean time to system failure can be represented as follows:

$$T = \int_0^{\infty} P(t) dt = \int_0^{\infty} [1 - (1 - e^{-\lambda t})^n] dt,$$

where $P(t)$ is the probability of no-failure.

Let us substitute variables

$$1 - e^{-\lambda t} = x \Rightarrow t = \frac{1}{\lambda} \ln \frac{1}{1-x} \Rightarrow dt = \frac{1}{\lambda(1-x)} dx.$$

Then,

$$T = \frac{1}{\lambda} \int_0^{\infty} \frac{1-x^n}{1-x} dx = \frac{1}{\lambda} \int_0^{\infty} \frac{-(x-1)(x^{n-1} + x^{n-2} + \dots + x + 1)}{1-x} dx = \frac{1}{\lambda} \left(\left. \frac{x^n}{n} \right|_0^1 + \left. \frac{x^{n-1}}{n-1} \right|_0^1 + \dots + \left. \frac{x^2}{2} \right|_0^1 + \left. x \right|_0^1 \right) = \frac{1}{\lambda} \left(1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n} \right).$$

Thus, for a “ m -out-of- n ” system we obtain:

$$T = \frac{1}{\lambda} \left(\frac{1}{n} + \frac{1}{n-1} + \frac{1}{n-2} + \dots + \frac{1}{m} \right). \quad (3)$$

The initial data for calculation are the values of failure flow $\lambda = 1.667 \times 10^{-6}$ 1/km and recovery flow $\mu = 0.041$ 1/km of one door.

By substituting the numerical values into (3), we obtain the following as summarized in Table 1.

Table 1. Results of failure flow calculation for the various block diagrams of a set of car doors

Type of diagram	Graphical representation	Failure flow parameter, λ , km^{-1}
Automatic side doors backed-up by manual doors		$\lambda_c = 5.5566 \times 10^{-6}$
An “ m -out-of- n ” system with sequentially connected gangway doors		$\lambda_c = 6.1917 \times 10^{-6}$
A “5-out-of-6” system that implies the option of manual operation for electromechanical doors		$\lambda_c = 7.88 \times 10^{-6}$
A “5-out-of-6” system that takes into account the possible manual operation of electromechanical side and gangway doors		$\lambda_c = 6.7694 \times 10^{-6}$

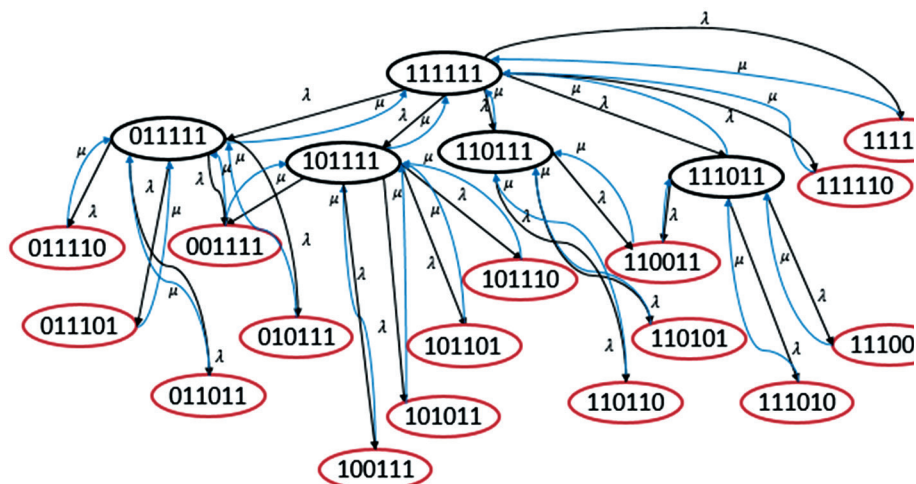


Fig. 6. State graph for the “3-out-of-4” structure with sequentially connected gangway doors

We obtained a series of failure flow values that depend on which of the reliability block diagrams can be considered. A team of engineers that was involved with this type of doors at various life cycle stages has adopted the “3-out-of-4” diagram (Fig. 5a).

Due to the fact that the logic-and-probability approach, although convenient for calculation, does not take into account the restoration of doors, the Markovian calculation of the failure flow was considered.

Calculation of the failure flow parameters of a set of car doors based on Markovian random processes

Let us estimate the dependability indicators of a set of car doors for the “3-out-of-4” structure with sequentially connected gangway doors using Markovian analysis. After determining the system state using the known *transition probability densities* λ and μ of the doors, let us construct a state graph for this structure (Fig. 6).

Next, let us construct a system of Kolmogorov differential equations, where the number of equations in the system will be equal to the number of states:

$$\begin{aligned} \frac{dP_{111111}}{dt} &= -6\lambda P_{111111}(t) + \mu [P_{011111}(t) + P_{101111}(t) + P_{110111}(t) + P_{111011}(t) + P_{111101}(t) + P_{111110}(t)]; \\ \frac{dP_{011111}}{dt} &= -(5\lambda + \mu)P_{011111}(t) + \mu [P_{001111}(t) + P_{010111}(t) + P_{011011}(t) + P_{011101}(t) + P_{011110}(t)] + \lambda P_{111111}(t); \\ \frac{dP_{101111}}{dt} &= -(5\lambda + \mu)P_{101111}(t) + \mu [P_{001111}(t) + P_{010111}(t) + P_{100111}(t) + P_{101011}(t) + P_{101101}(t)] + \lambda P_{111111}(t); \\ \frac{dP_{110111}}{dt} &= -(5\lambda + \mu)P_{110111}(t) + \mu [P_{001111}(t) + P_{010111}(t) + P_{100111}(t) + P_{101011}(t) + P_{110011}(t)] + \lambda P_{111111}(t); \\ \frac{dP_{111011}}{dt} &= -(5\lambda + \mu)P_{111011}(t) + \mu [P_{010111}(t) + P_{011011}(t) + P_{110011}(t) + P_{110101}(t) + P_{110110}(t)] + \lambda P_{111111}(t); \\ \frac{dP_{111101}}{dt} &= -\mu P_{111101}(t) + \lambda P_{111111}(t); \\ \frac{dP_{111110}}{dt} &= -\mu P_{111110}(t) + \lambda P_{111111}(t); \\ \frac{dP_{011110}}{dt} &= -\mu P_{011110}(t) + \lambda P_{011111}(t); \\ \frac{dP_{101110}}{dt} &= -\mu P_{101110}(t) + \lambda P_{101111}(t); \\ \frac{dP_{110110}}{dt} &= -\mu P_{110110}(t) + \lambda P_{110111}(t); \\ \frac{dP_{001111}}{dt} &= -2\mu P_{001111}(t) + \lambda [P_{011111}(t) + P_{101111}(t)]; \\ \frac{dP_{010111}}{dt} &= -2\mu P_{010111}(t) + \lambda [P_{011111}(t) + P_{101111}(t)]; \\ \frac{dP_{011011}}{dt} &= -2\mu P_{011011}(t) + \lambda [P_{011111}(t) + P_{101111}(t)]; \\ \frac{dP_{100111}}{dt} &= -2\mu P_{100111}(t) + \lambda [P_{011111}(t) + P_{101111}(t)]; \\ \frac{dP_{101011}}{dt} &= -2\mu P_{101011}(t) + \lambda [P_{011111}(t) + P_{101111}(t)]; \\ \frac{dP_{101101}}{dt} &= -\mu P_{101101}(t) + \lambda P_{101111}(t); \\ \frac{dP_{101110}}{dt} &= -\mu P_{101110}(t) + \lambda P_{101111}(t); \\ \frac{dP_{110011}}{dt} &= -2\mu P_{110011}(t) + \lambda [P_{101111}(t) + P_{110111}(t)]; \\ \frac{dP_{110101}}{dt} &= -\mu P_{110101}(t) + \lambda P_{110111}(t); \\ \frac{dP_{110110}}{dt} &= -\mu P_{110110}(t) + \lambda P_{110111}(t); \\ \frac{dP_{111001}}{dt} &= -\mu P_{111001}(t) + \lambda P_{111011}(t); \\ \frac{dP_{111010}}{dt} &= -\mu P_{111010}(t) + \lambda P_{111011}(t); \end{aligned}$$

Using MATLAB, let us evaluate the system's probability of no-failure. We will obtain $P = 0.9917$. This probability value is calculated for a period equal to the service life of the doors.

Let us use a simpler method of calculating the reliability indicators. For this purpose, let us use a Laplace transformation for a system of differential equations [8]:

$$\begin{cases} \frac{dP_{111111}}{dt} = -6\lambda P_{111111}(t) + \mu [P_{011111}(t) + P_{101111}(t) + P_{110111}(t) + P_{111011}(t) + P_{111101}(t) + P_{111110}(t)]; \\ \frac{dP_{011111}}{dt} = -(5\lambda + \mu)P_{011111}(t) + \lambda P_{111111}(t); \\ \frac{dP_{101111}}{dt} = -(5\lambda + \mu)P_{101111}(t) + \lambda P_{111111}(t); \\ \frac{dP_{110111}}{dt} = -(5\lambda + \mu)P_{110111}(t) + \lambda P_{111111}(t); \\ \frac{dP_{111011}}{dt} = -(5\lambda + \mu)P_{111011}(t) + \lambda P_{111111}(t). \end{cases}$$

We use Laplace transformations [8, 11] in order to obtain an expression for the mean time to failure, out of which we deduce the numerical value of the failure flow. For $P(t)$, we will obtain the following representation:

$$P(z) = \int_0^{\infty} P(t)e^{-zt} dt.$$

Then, we obtain the following system of equations:

$$\begin{cases} -1 = -6\lambda T_{111111} + \mu [T_{011111} + T_{101111} + T_{110111} + T_{111011}]; \\ 0 = -(5\lambda + \mu)T_{011111} + \lambda T_{111111}; \\ 0 = -(5\lambda + \mu)T_{101111} + \lambda T_{111111}; \\ 0 = -(5\lambda + \mu)T_{110111} + \lambda T_{111111}; \\ 0 = -(5\lambda + \mu)T_{111011} + \lambda T_{111111}. \end{cases}$$

By solving the system, we obtain:

$$\lambda_{\text{sys}} = 3.334 \times 10^{-6} \text{ 1/km if } \lambda = 1.667 \times 10^{-6} \text{ 1/km and } \mu = 0.041 \text{ 1/km.}$$

Thus, using Markovian analysis and Laplace transformation, the rated value of the failure flow parameter of a set of car doors was calculated based on the data on the failure and recovery flow parameters of a single door.

Conclusion

The paper elaborated upon a number of block diagrams for calculating the reliability indicators of sets of passenger car doors. Unlike the logic-and-probability approach to defining the failure flow parameter, the Markovian method is of higher priority, since it takes into account the recovery factor. Thus, the latter value of the system failure flow parameter is applicable to the “3-out-of-4” reliability block diagram under consideration for the purpose of monitoring the dependability indicators of sets of car doors in the course of operational testing.

The approaches examined in the paper enable the rolling stock component consumers and the manufacturer to agree on the methods of dependability calculation of

complex systems and to avoid conflicting methods of reliability indicator calculation when moving from an element to a system. The described method is currently used for calculating the rated value of the mean time to failure of a set of car doors installed in long-distance trains.

The research was performed with the financial support of RFBR as part of research project no. 20-31-70001.

References

1. Gnedenko B.V., Beliaev Yu.K., Soloviev A.D. [Mathematical methods in the dependability theory]. Moscow, Nauka; 1965. (in Russ.)
2. Shishmariov V.Yu. [Dependability of technical systems]. Moscow: Academia; 2010. (in Russ.)
3. Atovmian I.O., Vairadyan A.S., Rudnev Yu.P. et al. [Dependability of automated management systems]. Moscow: Vysshaya Shkola; 1979. (in Russ.)
4. GOST R 51901.14-2007. Risk management. Reliability block diagram and Boolean methods. Moscow: Standartinform; 2008. (in Russ.)
5. Viktorova V.S., Stepaniants A.S. [Models and methods of technical system dependability calculation]. Moscow: LENAND; 2014. (in Russ.)
6. Bertsche B. Reliability in Automotive and Mechanical Engineering. Springer; 2008.
7. Rausand M. Reliability of safety-critical systems: theory and application. Wiley; 2014.
8. Golinkevich T.A. [Applied dependability theory]. Moscow: Vysshaya Shkola; 1977. (in Russ.)
9. Polovko A.M. [Fundamentals of the dependability theory]. Moscow: Nauka; 1964. (in Russ.)
10. Druzhinin G.V. [Dependability of automation systems. Second edition]. Moscow: Energia; 1967.
11. Voronov A.A. [Automatic control theory]. Moscow: Vysshaya Shkola; 1986. (in Russ.)

About the authors

Maria V. Belousova, post-graduate student, Department of Modelling of Economic Systems, Saint Petersburg State University, Saint Petersburg, Russian Federation, e-mail: 27bmw1993@mail.ru.

Vitaly V. Bulatov, Candidate of Engineering, Senior Lecturer in Electromechanics and Robotics, Saint Petersburg State University of Aerospace Instrumentation, Saint Petersburg, Russian Federation, e-mail: bulatov-vitaly@yandex.ru.

Nikolay V. Smirnov, Doctor of Physics and Mathematics, Professor, Head of Department of Economic System Modelling, Saint Petersburg State University, Saint Petersburg, Russian Federation, e-mail: n.v.smirnov@spbu.ru.

The authors' contribution

Belousova M.V. Calculated the parameters of the failure flow for block diagrams of a set of car doors, developed state graphs for the 3-out-of-4 set structure, evaluated the dependability indicators for a set based on Markovian processes.

Bulatov V.V. Described the object of research, generally evaluated the applicability of the considered mathematical methods to the given rolling stock systems and developed reliability block diagrams for a set of passenger car doors.

Smirnov N.V. Together with Belousova M.V. formulated the mathematical problem, verified the correctness of all arguments and calculations, proposed an interpretation of the obtained results.

Conflict of interests

The author declares the absence of a conflict of interests.

Simulation of railway marshalling yards using the methods of the queueing theory

Maksim L. Zharkov^{1*}, Mikhail M. Pavidis²

¹Matrosov Institute for System Dynamics and Control Theory, Siberian Branch of the Russian Academy of Sciences, Irkutsk, Russian Federation, ²Irkutsk State Transport University, Irkutsk, Russian Federation

*zharkm@mail.ru



Maksim L. Zharkov



Mikhail M. Pavidis

Abstract. Aim. The paper primarily aims to simulate the operation of railway transportation systems using the queueing theory with the case study of marshalling yards. The goals also include the development of the methods and tools of mathematical simulation and queueing theory. **Methods.** One of the pressing matters of modern science is the development of methods of mathematical simulation of transportation systems for the purpose of analysing the efficiency, stability and dependability of their operation while taking into account random factors. Research has shown that the use of the most mature class of such models, the single-phase Markovian queueing systems, does not enable an adequate description of transportation facilities and systems, particularly in railway transportation. For that reason, this paper suggests more complex mathematical models in the form of queueing networks, i.e., multiple interconnected queueing systems, where arrivals are serviced. The graph of a queueing network does not have to be connected and circuit-free (a tree), which allows simulating transportation systems with random structures that are specified in table form as a so-called “routing matrix”. We suggest using the BMAP model for the purpose of describing incoming traffic flows. The Branch Markovian Arrival Process is a Poisson process with batch arrivals. It allows combining several different arrivals into a single structure, which, in turn, significantly increases the simulation adequacy. The complex structure of the designed model does not allow studying it analytically. Therefore, based on the mathematical description, a simulation model was developed and implemented in the form of software. **Results.** The developed models and algorithms were evaluated using the case study of the largest Russian marshalling yard. A computational experiment was performed and produced substantial recommendations. Another important result of the research is that significant progress was made in the development of a single method of mathematical and computer simulation of transportation hubs based on the queueing theory. That is the strategic goal of the conducted research that aims to improve the accuracy and adequacy of simulation compared to the known methods, as well as should allow extending the capabilities and applicability of the model-based approach. **Conclusions.** The proposed model-based approach proved to be a rather efficient tool that allows studying the operation of railway marshalling yards under various parameters of arrivals and different capacity of the yards. It is unlikely to completely replace the conventional methods of researching the operation of railway stations based on detailed descriptions. However, the study shows that it is quite usable as a primary analysis tool that does not require significant efforts and detailed statistics.

Keywords: transportation, marshalling yard, mathematical simulation, queueing theory, queueing network, simulation, computational experiment, dependability.

For citation: Zharkov M.L., Pavidis M.M. Simulation of railway marshalling yards using the methods of the queueing theory. *Dependability* 2021;3: 27-34. <https://doi.org/10.21683/1729-2646-2021-21-3-27-34>

Received on: 27.01.2021 / **Upon revision:** 09.02.2021 / **For printing:** 17.09.2021.

Introduction

In recent years, the application of the queueing theory [1] for simulating transportation systems of various levels became an important and relevant area of research, as it allows assessing the efficiency, stability and operational dependability of transportation while taking into account random factors. At the same time, the conventional toolkit of single-phase Markovian queueing systems (QSs) proved to be unsatisfactory, as: a) with a few exceptions, in transportation systems, service is carried out in several stages (phases); b) incoming vehicles cannot be considered as individual arrivals, as they may have a complex structure and differ in terms of capacity. For instance, such are the railway trains arriving to a freight or marshalling station (MS). The number of cars in a train may differ; while their types are also typically different [2]. Thus, a more complex simulation approach is required that uses non-Markovian and/or multi-phase QSs, as well as queueing networks (QNs). At the same time, since such mathematical objects are difficult to analyse, algorithmic and software tools need to be developed that would enable computer simulation.

Source overview

The applicability of two-phase QSs for railway station simulation was mentioned as early as in [2], yet at that time no systematic studies were carried out in this area. The reason probably consisted in the insufficiently mature mathematics. In the 21-st century, the situation changed. Thus, the Irkutsk School lead by Academy Member I.V. Bychkov and RAS Prof. A.L. Kazakov, to which the authors belong, for more than 10 years has been developing [3] an area of research associated with the application of multi-phase QSs as a model for processing the arrivals to transportation nodes. The nature of the latter may vary greatly from a metropolitan transport hub [4] to a railway freight station [3]. Similar research is also conducted abroad [5-7]. In [5], the queueing theory (QT) is used for identifying the capacity of railway lines, in [6], it is used for the mathematical description of the operation of stations and infrastructure facilities, in [7], it helps simulate railway node operations. However, the QT is much more often applied to the information and telecommunications technologies.

There are a large number of schools of thought and groups of researchers involved in this area of research. Let us mention three of them, who, in the authors' opinion, occupy leading positions in all of the ex-Soviet countries: the Moscow school lead by Prof. V.V. Rykov (see, for example, [8-10]), the Tomsk school lead by Prof. A.A. Nazarov (see, for example, [11-13]) and the Belarusian school lead by Prof. A.N. Dudin (see, for example, [14-16]). Naturally, the list can be continued, yet this paper does not aim and cannot make a comprehensive review of the research findings in the field of QT application

for the purpose of simulating information systems and technologies.

Going back to the activities of the Irkutsk school, we should note that, as the simulation results show, the common features inherent to all transportation systems in this case prevail over the various differences and allow examining them together. Although, of course, any simulation approach requires an adaptation to the object of research, which allows taking into account the structure and directions of the internal traffic flows.

The incoming traffic flows [17] deserve a separate discussion. As a model that allows capturing their complex and heterogeneous structure, we propose using the *BMAP* (Branch Markovian Arrival Process) that enables an integration of a number of different arrivals [14]. Essentially, this is a generalized case of the Poisson stream with grouped arrivals. This model was first suggested by the Italian mathematician D. Lucantoni back in 1991 [18], yet until now it has only been used for information system simulation [16].

Based on three-phase QSs, the authors, along with the above-mentioned transport hubs (in Moscow and Ekaterinburg) and freight stations [3], constructed mathematical models of MS operations [19-21]. The results were tested with specific transportation facilities both in Russia and abroad, and attracted the interest of transportation experts. However, the proposed approach has also shown a weakness that is due to the fact that the apparatus of a multi-phase QS is only able to describe linearly structured systems. They do not support the organization of loop motion of arrivals, which, in particular, is typical for some MSs.

Methods

For the purpose of solving the above problems and extending the capabilities of the simulation approach, we propose using a new class of objects, the queueing networks [21, 22]. A QN is understood as a set of interconnected QSs within which arrivals circulate (are serviced) [22, 23]. Unlike in the multi-phase QSs, a QN graph does not have to be connected and circuit-free (a tree), which allows a much more flexible simulation of structurally complex transportation systems that are defined by a "routing matrix". Unfortunately, in this case, the functional extension of the simulation apparatus has a downside. The object of research becomes fundamentally more complex. If with the multi-phase QSs, analytical results can sometimes be obtained, for QNs, only a numerical study is available using single toss-based simulation methods [24] (the Monte Carlo methods).

This paper is a follow-up to [21]. Based on the previously proposed concept of railway transportation system simulation, the authors suggest a method of QN-based MS simulation, thus developing upon the previously obtained results in the QS-based simulation of the above transportation facilities [19]. In particular, we have developed the

appropriate numerical algorithms that are implemented as software that allows researching the properties and evaluating the parameters of multi-phase systems and queueing networks by means of specially organized computer experiments based on methods of statistical simulation. For the purpose of testing the proposed approach, a model station is examined. A computational experiment was carried out, conclusions were made regarding the specificity of the station's operation.

Mathematical model

As it is known, a marshalling system is a complex structure. It is designed for mass breaking-up of freight trains into individual groups of cars, their handling and accumulation with subsequent making of new trains out of them. The MSs perform standard operations and consist of similar elements. Let us identify the most important of them that will be taken into account in the mathematical model. Standard MS perform the following actions: acceptance of trains into the receiving yard (RY) and uncoupling of the locomotive; breaking-up of train on the hump; accumulation of cars in the marshalling yard (MY) in accordance with the train make-up plan; delivery of a train into the departure yard (DY), coupling of the locomotive and departure of the made train from the system. In each yard and at the hump there are service facilities of different capacities. The incoming train flow includes transit, local and other categories of trains, whose parameters may vary greatly. The trains arrive from several directions (two or more). An incoming train should be regarded as a group of arrivals, as cars are serviced independently from each other and occupy certain positions on the tracks of the yards. Therefore, the total incoming train flow consists of at least four sub-flows, each of which is a group of arrivals. Passenger trains usually bypass MYs and are not taken into account.

A significant part of the incoming train flow is transit trains that travel across the territory of Russia. This group is significantly affected by random factors due to the very long travel distances, therefore the traffic management system cannot effectively schedule all categories of trains on an individual railway line. As a consequence, trains significantly deviate from the schedule [25]. Therefore, it can be assumed that the arrival of trains is a random value.

The mathematical model of an MS is constructed in two stages. At the first stage, the incoming arrivals are described. For that purpose, a *BMAP* model is used that allows aggregating several different arrival flows into a single structure. At the second stage, the processing of arrivals in the system is described. In order to take into account the complex hierarchical structure of the system, it is suggested to use a QN.

The **BMAP arrival** (Batch Markovian Arrival Process) differs from the simple batch arrival in that: a) the rate of batch arrivals λ_v depends on the state number of the Markov

control chain v_t with continuous time and finite space-state $\{0, 1, \dots, W\}$; b) the time of Markov chain v_t being in state v has an exponential distribution with parameter λ_v ; c) after the time of the chain's being in state v has elapsed, it, with a given probability $p_k(v, v')$, will change into a different state v' ; a batch of size $k \geq 0$ is generated in the process; d) transition probabilities $p_k(v, v')$ comply with the normalization

requirement $\sum_{k=0}^{\infty} \sum_{v=0}^W p_k(v, v') = 1$. The transition rates of the Markov chain are conveniently stored in matrices

$$(D_0)_{v,v} = -\lambda_v, v = \overline{0, W}; (D_0)_{v,v'} = \lambda_v p_0(v, v'), v, v' = \overline{0, W}; (D_k)_{v,v'} = \lambda_v p_k(v, v'), v, v' = \overline{0, W}, k \geq 1. \quad (1)$$

A QN is the sum of the finite number S of QNs (herein after referred to as nodes), in which arrivals are transferred from one node to another in accordance with the routing matrix P [22, 23]. Let us accept that arrivals are received into a QN from an external source. If it is accepted as an additional node with index 0, the route of an arrival is defined by the stochastic matrix $P = \|P_{ij}\|$ of size $(S+1) \times (S+1)$. Its elements are P_{ij} , the probabilities of the arrival moving from node i to node j ($i, j = \overline{1, S}$), P_{0j} and P_{j0} , are, respectively, the probability of an arrival into the j -th node from the source and the probability of an arrival leaving the network after being serviced in j -th node ($j = \overline{1, S}$). Obviously, $\sum_{j=0}^S P_{ij} = 1$ ($i = \overline{0, S}$), $P_{00} = 0$ [22, 23].

Marshalling yard simulation

Let us examine a model marshalling station (MMS). Its characteristics correspond to the Ekaterinburg-Sortirovochny (E-S) station of the Sverdlovsk Railway, the largest MS in Russia. E-S is a two-system MS with a serial yard arrangement that handles trains from five lines: 1) Tagil, 2) Kungur, 3) Kazan, from the stations, 4) Ekaterinburg-Tovarny and 5) Ekaterinburg-Passazhirsky. The down system serves lines 4) and 5), the up system serves lines 1), 2), 3). The two systems are almost identical in terms of the performed operations. The up system is currently undergoing an upgrade and we are unaware of its performance parameters. Therefore, the MMS uses the up system numbers (see Fig. 1): the RY includes 11 specialized tracks for receiving freight trains with a total capacity of 716 conventional cars (conv. cars) serviced by two humping engines that move trains to the hump along two tracks with a total capacity of about 100 conv. cars; the hump has a large capacity and can handle up to 5500 conv. cars per day; the MY has 35 tracks with the total capacity of 2535 conv. cars, train formation is ensured by three engines that move trains to the DY; it has 15 tracks with the total capacity of 980 conv. cars, trains depart from it along three lines. The tracks in the yards differ in length, the longest being able to accommodate over 80 conv. cars. E-S accepts a car flow that includes three categories of trains: a) transit with rehandling; b) transit without rehandling; c) local. The a) and c) trains arrive to the receiving yard and undergo the

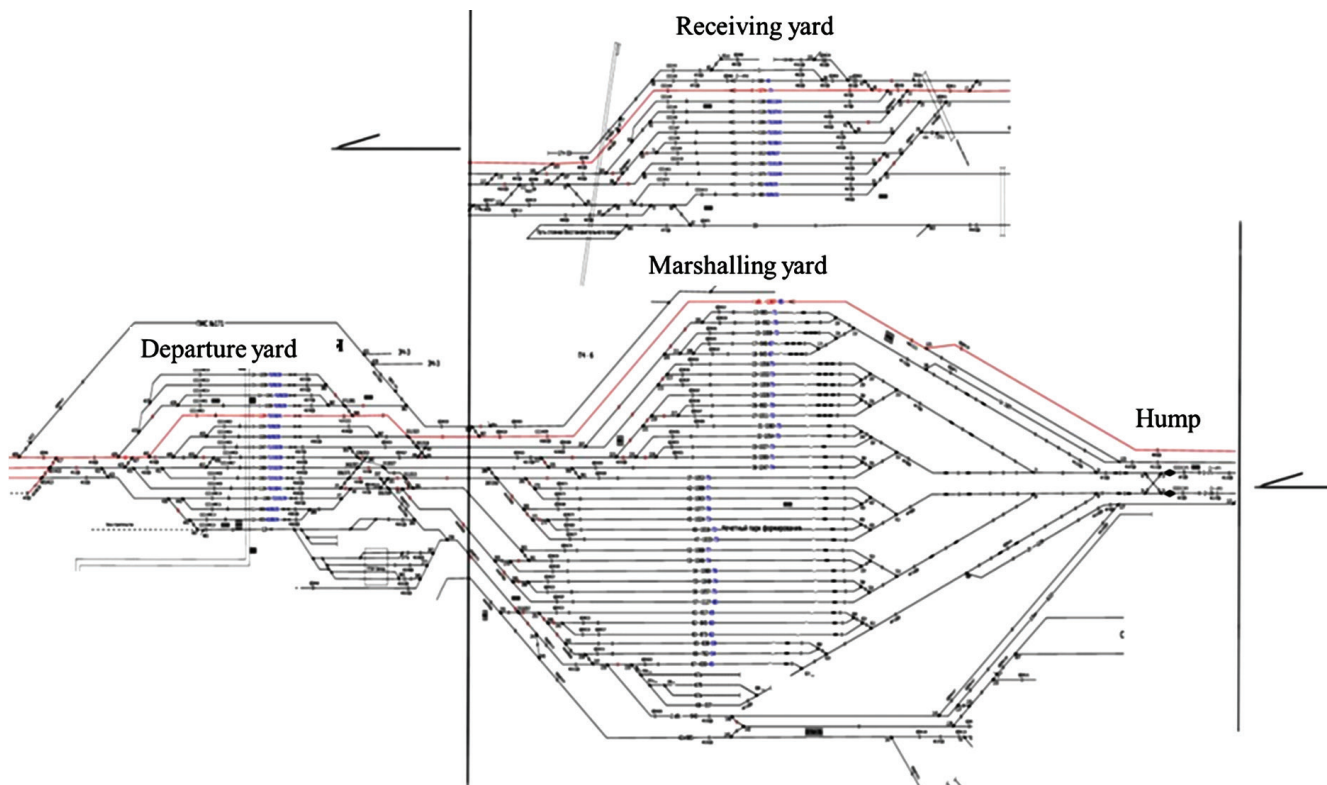


Fig. 1. Down system of the E-S marshalling station

rest of the service. The b) trains are serviced in the RY or DY, then leave the system.

While preparing paper [25], we collected statistical data on trains operating on the Sverdlovsk Railway, in particular, freight trains: actual and scheduled times of arrival at station, number of cars in trains and their numbers. They might be partially obsolete, but we did not manage to obtain more recent data, therefore further station simulation will be based on the available information. They show that in a half of the cases, schedule violations by freight trains exceed 30 minutes and more than two hours in 28% of cases. Consequently, the time between the arrival of trains to the station can be taken as a random value. The number of cars in the transit and most local freight trains follows a binomial distribution $B(80, 0.9)$. Their average number is 70, the maximum number is 80 for all categories.

In the MY, the hump is considered to be the primary facility that defines the system's performance. Its capacity defines the size of the car stream the station can handle. We were unable to obtain the statistics of the incoming train traffic, so we shall calculate the number of the trains arriving for breaking-up in such a way as to make the hump loading 70% of the maximum possible. This value corresponds to the average planned loading for

such facilities. Then, the system is to accept 3850 cars or 55 trains per day. We shall take the number of trains without breaking-up to be one quarter of the number of those broken-up, i.e., 14 per day. We deduce that the MMS accepts 69 freight trains from two lines per day, i.e., 35 trains from line 4) and 34 from line 5). In the model, we divide the trains depending on the handling procedure at the station, i.e., with breaking-up (trains A) and with no breaking-up (trains B), and the line. Thus, the incoming car flow will consist of four sub-flow groups. We will use a *BMAP* flow model for the purpose of mathematical description of such car flow. It will include $81 \times 4 D_k$ matrices, $k = \overline{0, 80}$. Their elements are calculated using formulas (1), where $\lambda_0 = \lambda_1 = \lambda_2 = \lambda_3 = \lambda \cdot 69 / 24 = 2.875$, $p_0 = 28 / 69 = 0.41$, $p_2 = 27 / 69 = 0.39$, $p_1 = p_3 = 7 / 69 = 0.1$, $p_k(v, v') = p_v \cdot f(k)$, $v, v' = \overline{0, 3}$, $k = \overline{0, 80}$, $f(k)$ is the probability of the arrival of a group of cars of size k that follows the law $B(80, 0.9)$.

The MMS model in the form of a QN is as follows. The system has four service nodes: 1) in the RY, the two shunting engines are the channels, the yard tracks are the queue, then Node 1 is a QS with two channels and a queue with 716 positions; 2) at the hump, one hump track and the shunting device are one channel, the second track is the queue, then Node 2 is a QS with one channel and a

Table 1. Parameters of the channel operation in the MMY subsystems

	Node 1	Node 2	Node 3	Node 4
F	$N(20, 2)$	$N(20, 3)$	$N(40, 5)$	$N(40, 3)$
X	$B(80, 0.9)$	$B(80, 0.9)$	$B(80, 0.9)$	$B(80, 0.9)$

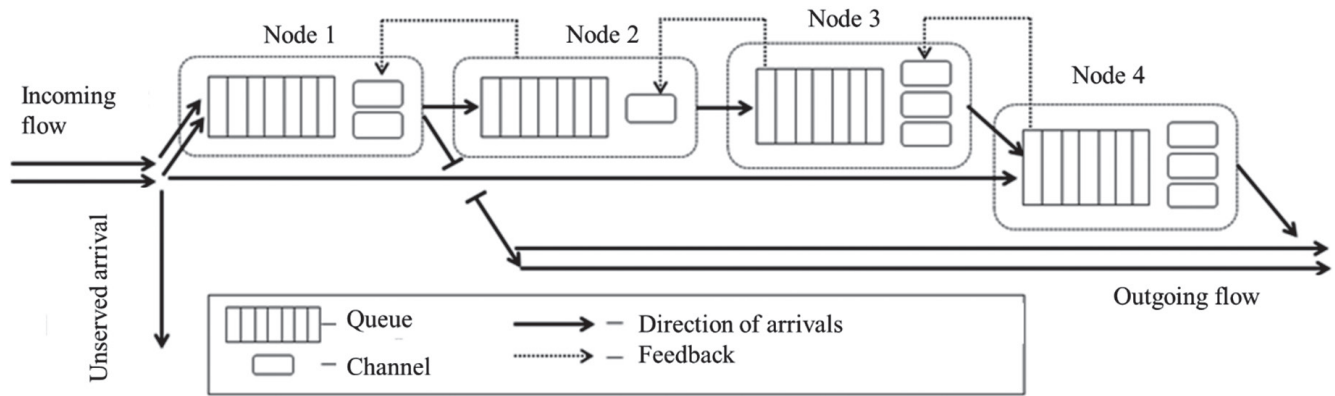


Fig. 2. QN diagram

queue with 100 positions; 3) in the MY, the three engines are the channels, the yard tracks are the queue, then Node 3 is a three-channel QS with a queue with 2535 positions; 4) in the DY, the three departure tracks are the channels (three lines), the yard tracks are the queue, then Node 4 is a three-channel QS with 980 positions. The service time in the node channels is taken as a random normally distributed value. The choice is due to the fact that the station staff strive to bring the duration of technical operations closer to the standard values. Deviations are due to the effect of random factors. An example would be a train that has cars with only two or three lines (destination stations). Then, it can be broken-up on the hump in 12 minutes instead of the standard 20. The MMS operation parameters are presented in Table 1.

The channel performance parameters (see Table 1) are defined based on the standard MY operation process taking

into account the specificity of E-S. The service time in Node 1 is taken as the time of humping from the departure park and return of the engine. The service time interval is [12.0; 28.0] minutes, the average time is 20 minutes. The hump capacity (Node 2) is 5500 conv. cars per day or 4 conv. cars per minute. Let the service time interval be [11.0; 29.0] minutes. In Node 3 (MY), the service time is taken as the time of engine coupling to the train, completion of train formation and its relocation to RY, as well as the return of the engine to the MY. The service time interval is [25.0; 55.0] minutes, the average time is 40 minutes. In the DY (Node 4) we take into account the time of engine coupling to the train, short brake testing and the train's departure from the system. Let the average service time be 40 minutes and the interval be [31.0; 49.0] minutes.

Trains A and B arrive into the system. The former enter Node 1 and proceed along the entire system. The latter are

Table 2. Results of experiment 1

	Received	Lost	$T_{\text{sys}} (m)$	P_G	P_R
Groups of arrivals	2431.80	0	161.76	0	0
Arrivals	175199.20	0			
	K	l	$t_{ph} (m)$	$t_{lock} (m)$	P_{lock}
Node 1	1.00	41.46	35.68	2897.00	0.05
Node 2	0.74	24.77	32.61	0.00	0.00
Node 3	1.43	0.00	42.43	0.00	0.00
Node 4	1.59	26.55	51.05	-	-

Table 3. Results of experiment 2

	Received	Lost	$T_{\text{sys}} (m)$	P_G	P_R
Groups of arrivals	2903.50	0	166.13	0	0
Arrivals	209049.50	0			
	K	l	$t_{ph} (m)$	$t_{lock} (m)$	P_{lock}
Node 1	1.23	61.10	39.20	4681.25	0.08
Node 2	0.83	32.65	34.48	0.00	0.00
Node 3	1.60	0.00	42.44	0.00	0.00
Node 4	1.79	26.64	50.02	-	-

Table 4. Results of experiment 3

	Received	Lost	$T_{\text{sys}} \text{ (m)}$	P_G	P_R
Groups of arrivals	3453.50	1.00	178.23	0.0003	0.0003
Arrivals	248605.50	72.82			
	K	l	$t_{ph} \text{ (m)}$	$t_{\text{lock}} \text{ (m)}$	P_{lock}
Node 1	1.50	111.15	48.60	7216.50	0.12
Node 2	0.90	42.58	36.96	0.00	0.00
Node 3	1.74	0.00	42.45	0.00	0.00
Node 4	1.97	29.85	50.23	-	-

received in Node 1 or Node 4, are serviced and leave the system. We use a routing matrix to account for different train service procedures. Let us assume that trains B may be admitted to Node 1 or Node 4 with an equal probability of $p_b = 0.5$. Then, trains from the total incoming flow arrive at Node 4 with the probability of $p_{0,4} = (2 \cdot 7 / 69)p_b = 0.1$, and at Node 1 with the probability of $p_{0,1} = 1 - p_{0,4} = 0.9$. After servicing in Node 1, trains B leave the system with the probability of $p_{1,0} = 7 / 62 = 0.11$. Trains A go to Node 2 with the probability of $p_{1,2} = 1 - p_{1,0} = 0.89$, then proceed across Nodes 3 and 4. In QT terms, the MMY model will be of the form of QN with routing matrix P formed by the above probabilities and the following nodes: 0 is the source of the incoming flow; 1 is $BMAP / G^B / 2 / 716$; 2 is $* / G^B / 1 / 100$; 3 is $* / G^B / 3 / 2535$; 4 is $* / G^B / 3 / 980$, where B is the binomial distribution, G is the random law of service time distribution. The diagram of the above QN is shown in Fig. 2.

Let us study the obtained QN numerically using the simulation model [24]. The sought characteristics, the performance indicators, are [1, 22, 23]: P_R and P_G are the probabilities of handling an application and a group of applications, T_{sys} is the application's average system time, l is the average queue length, K is the average number of busy channels, t_{ph} is an application's average time in the node, t_{lock} is the total blocking time of all of a node's channels, P_{lock} is the probability of blocking of a single channel of a node.

Computational experiment. Tables 2-4 below show the results of scenario-based simulation of the operation of the above QN (see Fig. 2) under various parameters of the incoming car flow. Each table shows the average results of 10 program starts. The simulation time is five weeks for all experiments, the minimal time required for the simulation model to calculate the stationary QN characteristics. The primary indicator of the fact the QN is successfully handling the load is $P_R = 0$, which, for large industrial and transportation systems means the absence of the risk of accidents. Thus, in particular, this indicates that the station handles all trains and the operation is fault-free.

Experiment 1. Table 2 shows simulation results for the case of car flow $\lambda = 4$ per hour.

Experiment 2. Table 3 shows the simulation results for the case of total flow of 5880 cars or almost 84 trains per day, of which 17 do not require handling ($\lambda = 3.5$).

Experiment 3. Table 4 shows simulation results for the case of arrival rate $\lambda = 4$ per hour.

Discussion of the simulation results

Out of the results of Experiment 1, it can be seen that the average system dwell time of an arrival (car) is 2.7 hours. The probability of failure is zero, i.e., the car flow is accepted to the station continuously, which is an important indicator of the dependability of railway transportation. Thus, the MMY successfully handles the car flow; a dependable and reliable station operation is ensured.

Experiment 2. Node 1 is the busiest as the average time of channel blocking is 2.2 hours a day, which is not critical, yet it affects the station's operation. In such situation, the traffic controllers arrange freight trains at adjacent railway stations, where they await the release of tracks in the receiving yard. The marshalling station operates normally and handles the load, however, due to longer blocking of Node 1 channels, the goods delivery time increases, while the stability of the transportation system is not fully ensured.

Experiment 3. The results of the simulation show that the channels of Node 1 are on average blocked for 3.4 hours a day, which is critical for the system, since the queue of Node 1 overflows and the probability of failure becomes non-zero, dependable and reliable operation of MMY is not ensured.

Based on the results of all experiments, the following general conclusion can be made. The system operates normally at maximum rate of incoming train flow of $\lambda = 3.5$ trains per hour. The limit value of a station's operation is the rate of incoming train flow of $\lambda = 4$ trains per hour. The system's bottleneck is the insufficient capacity of the receiving yard (Node 1). In order to relieve the load, all trains should be redirected to Node 4 (DY) without handling. Further improvement of the system's performance requires increased hump capacity. However, that will require its reconstruction involving substantial material costs.

Conclusion

Summing up the results of the research, we should note that, as the authors hope, it is only the first major step on a long way that is to culminate in the creation

of a single methodology for mathematical and computer simulation of transportation hubs using QN. That will enable improved accuracy and adequacy of the simulation aimed at assessing the efficiency, stability and operational dependability of transportation systems, as well as enhancing the model-based approach. Another important problem to be solved as part of the transportation systems simulation is that the software products are alienable from the developers. The most natural solution would be to use intellectualization tools in the creation of an intelligent system for managing the development of the region's transportation and logistics infrastructure [26] led by academy member I.V. Bychkov.

In conclusion, let us note that it is unlikely that the proposed model-based approach will completely replace the conventional methods of studying the operation of railway stations based on their detailed description. However, as we have identified, it can quite be used as a primary analysis tool that does not require significant efforts and detailed statistics.

Acknowledgements

The authors express their gratitude to Professor A.L. Kazakov for his useful advice in the preparation of the paper's materials and assistance in its writing. The research was carried out with the financial support of RFBR as part of research project no. 20-010-00724; RFBR and the Government of the Irkutsk Oblast as part of research project no. 20-47-383002.

References

1. Gnedenko B.V., Kovalenko B.V. [Introduction into the queueing theory]. Moscow: LKI; 2007. (in Russ.)
2. Akulinichev V.M., Kudriavtsev V.A., Koreshkov A.N. [Mathematical methods in the operation of railways]. Moscow: Transport; 1981.
3. Kazakov A.L., Maslov A.M. [Construction of a model of an uneven traffic flow. Case study of a railway freight station]. *Modern Technologies. Systems Analysis. Modeling* 2009;3:27-32. (in Russ.)
4. Zhuravskaya M.A., Kazakov A.L., Zharkov M.L. et al. Simulating the operation of transport hub of a metropolis as a three-phase queueing system. *Transport of the Urals* 2015;3:17-22. (in Russ.)
5. Weik N., Niebel N. Capacity analysis of railway lines in Germany – a rigorous discussion of the queueing based approach. *Journal of Rail Transport Planning & Management* 2016;6(2):99-115.
6. Dorda M., Teichmann D. Modelling of freight trains classification using queueing system subject to breakdowns. *Mathematical Problems in Engineering* 2013;2013:11.
7. Nießen N. Waiting and loss probabilities for route nodes. In: *Proceedings of the 5th International Seminar on Railway Operations Modelling and Analysis*. Copenhagen (Denmark); 2013.
8. Bronshtein O.I., Raykin A.A., Rykov V.V. [On a queueing system with losses]. *Izvestiia Akademii nauk SSSR: Tekhnicheskaya kibernetika* 1964;4:39-47. (in Russ.)
9. Kitaev M.Yu., Rykov V.V. Controlled queueing system. N.Y.: CC Press; 1995.
10. Rykov V.V. Controllable queueing systems from the very beginning up to nowadays. *Materials of Information Technologies and Mathematical Modelling named after A.F. Terpugov* 2017;1:25-26.
11. Grachev V.V., Moiseev A.N., Nazarov A.A. et al. Multistage queueing model of the distributed data processing system. *Proceedings of TUSUR University* 2012;2-2(26):248-251. (in Russ.)
12. Shklennik M., Moiseeva S., Moiseev A. Optimization of two-level discount values using Queueing tandem model with feedback. *Communications in Computer and Information Science* 2018;912:321-332.
13. Galileyskaya A.A., Lisovskaya E.Yu., Moiseeva S.P. et al. On sequential data processing model that implements the backup storage. *Modern Information Technologies and IT-Education* 2019;3:579-587. (in Russ.)
14. Dudin A.N., Klimenok V.I. [Queueing systems with correlated flows]. Minsk: BSU; 2000. (in Russ.)
15. Kim C., Dudin A., Dudina O. et al. Tandem queueing system with infinite and finite intermediate buffers and generalized phase-type service time distribution. *European Journal of Operational Research* 2014;23:170-179.
16. Vishnevskii V.M., Dudin A.N. Queueing systems with correlated arrival flows and their applications to modeling telecommunication networks. *Automation and remote control* 2017;8:3-59. (in Russ.)
17. Gasnikov A.V., editor. [Introduction into the mathematical modelling of traffic flows]. Moscow: MIPT; 2010. (in Russ.)
18. Lucantoni D.M. New results on single server queue with a batch Markovian arrival process. *Commun. Statist. Stochastic Models* 1991;7:1-46.
19. Kazakov A.L., Pavidis M., Zharkov M.L. Multiphase systems of mass service in switchyard modeling. *Herald of the Ural State University of Railway Transport* 2018;2:4-14.
20. Kazakov A.L., Pavidis M. On one approach to model operation of marshalling stations. *Transport of the Urals* 2019;1(60):29-35. (in Russ.)
21. Zharkov M.L., Pavidis M.M. Modelling of railway stations based on queueing networks. *Aktualnye Problemy Nauki Pribaikalia* 2020;3:79-84. (in Russ.)
22. Walrand J. An introduction to queueing networks. Mir; 1993.
23. Ivnitky V.A. Queueing network theory. Moscow: FIZMATLIT; 2004. (in Russ.)
24. Kelton D.W., Law A.M. Simulation modelling and analysis. Saint Petersburg: Piter; 2004.
25. Zharkov M.L., Parsyurova P.A., Kazakov A.L. Modeling operation of railway stations and rail network sections based on studying train schedule deviations.

Proceedings of Irkutsk State Technical University 2014;6(89):23-31. (in Russ.)

26. Bychkov I.V., Kazakov A.L., Lempert A.A. et al. [An intelligent system for managing the transportation and logistics infrastructure of a region]. *Control Sciences* 2014;1:27-35. (in Russ.)

About the authors

Maxim L. Zharkov, Candidate of Engineering, Researcher, Matrosov Institute for System Dynamics and Control Theory, SB RAS, Shelekhov, Irkutsk Oblast, Russian Federation, e-mail: zharkm@mail.ru

Mikhail M. Pavidis, post-graduate student, Irkutsk State Transport University, Irkutsk, Russian Federation, e-mail: pavidismiha1994@mail.ru

The authors' contribution

Zharkov M.L. stated the problem and developed the model-based approach. **Pavidis M.M.** made a source review, performed model identification and computational experiment, drew conclusions based on the simulation results.

Conflict of interests

The authors declare the absence of a conflict of interests.

Use of the terms “estimate” and “definition” in dependability-related standards

Boris P. Zelentsov, Siberian State University of Telecommunications and Information Sciences, Novosibirsk, Russian Federation,
zelentsovb@mail.ru



Boris P. Zelentsov

Abstract. The paper aims to improve the terminology used in dependability-related state standards. Examples are given of the use of the terms “estimate” and “definition” in the “Risk management” and “Dependability in technics” series of state standards. The meanings of those terms were clarified based on the existing regulatory documents. Requirements for the integrity of the used terms were defined. Wordings were proposed for the term definitions that feature the words “estimate” and “definition”. **Aim.** To examine and discuss the common, but not sufficiently substantiated terms “estimate” and “definition” used in state standards, i.e., to consider the legitimacy of their application as part of the above series of state standards. Proposals as to the improvement of such terms’ application were also set forth. **Methods.** Examples are given of the use of the terms “estimate” and “definition” in state standards. Based on the existing state standards, the actual meanings of the considered terms were clarified: “definition” refers to the way a term is defined, while “estimate” and “estimation” are closely associated with mathematical statistics. The requirements for the integrity of the used terminology are defined and come down to it being unambiguous, consistent within itself and across the relevant state standards. In this context, the shortcomings of the examined terms are shown that are associated with the above requirements, i.e., the meaning, content, essence and key features of such terms are clearly defined. Any comments or references to other regulatory documents are missing as well. **Results.** In most standards, in the “Terms and definitions” section, the concept of “definition” is used correctly, i.e., terms are defined. However, in other cases, the concept of “definition” is used in a different sense, as nothing is actually being defined. Based on the term integrity requirements and in light of the above shortcomings, proposed replacements for the terms in question were defined. In most cases, instead of the terms “estimate” and “definition”, it is proposed to use the terms “calculation” and “computation”, as well as their cognates, “calculate”, “compute”. It should be noted that along the state standards, these terms are used in technical documentation, science papers, monographs and textbooks. **Conclusions.** The use of the examined terms in some standards lacks integrity. The requirements of the standardization recommendations are not observed, the terms are not unambiguous and consistent with other standards. Based on these requirements, the paper proposes improved ways of using the terms “estimate” and “definition”. The suggested terms should be considered as a tentative proposal. Final definitions and/or replacements of these terms are to be developed through extensive discussion and compromise.

Keywords: dependability, dependability-related terminology.

For citation: Zelentsov B.P. Use of the terms “estimate” and “definition” in dependability-related standards. *Dependability* 2021;3: 35-38. <https://doi.org/10.21683/1729-2646-2021-21-3-35-38>

Received on: 26.02.2021 / **Upon revision:** 22.07.2021 / **For printing:** 17.09.2021.

Introduction

This paper examines the terms used in the “Risk management” and “Dependability in technics” series of state standards. The terms “estimate” and “definition” are used in those standards.

This paper aims to examine the applicability of these terms in the context of specific standards.

The author justifies the use of the terms “estimate” and “definition” relying on the existing state standards [2, 3, 12]. Standard [3] has been replaced with standard [2]. However, in the course of the development and application of the standards of the above series up to 2019, standard [3] was in force. Therefore, the paper refers to both standards of the “Statistical methods” series, the more so since they do not significantly differ in terms of the definitions of primary terms.

Source overview

Most state standards have a “Terms and definitions” section. In particular, those sections feature:

- 16 terms in [1] (GOST R 27.302-2009);
- 6 terms in [4] (GOST R 51901.5-2005);
- 8 terms in [5] (GOST R 51901.12-2007);
- 15 terms in [9] (GOST R IEC 61165-2019).

If the section is blank, a reference to another state standard is provided.

The term “estimate” is commonly used in state standards. Let us note examples of this term’s application:

- In [1] (GOST R 27.302-2009):
 - reliability estimate (6.1.3, 6.13);
 - probability estimate (7.5.4.3);
 - gate estimate (7.5.2.2);
 - fault tree estimate (6.1.2, 7);
 - estimating the probability of failure (5.4.3).
- In [4] (GOST R 51901.5-2005):
 - dependability estimate (A.1.12.1);
 - dependability indicators estimate (4.1);
 - dependability improvement estimate (4.1);
 - estimate of primary event characteristics (4.4.1).
- In [5] (GOST R 51901.12-2007):
 - estimate of the probability of failure (5.2.9);
 - estimate of failure rate (5.3.4);
 - calculated or estimated probability of failure (5.3.6.2).

In [7] (GOST R 51901-14-2007):

- estimate of the probability of no-failure (7.2, 8.3);
- probability of no-failure can be estimated using the formula (8.1.1).

- In [9] (GOST R IEC 61165-2019):
 - estimate of the probability of no-failure (B.1).

Let us note the cases of the application of the term “definition”:

- In [1] (GOST R 27-302-2009):
 - probabilities are defined in the usual way (6.1.3).
- In [4] (GOST R 51901-5-2005):

- definition of numerical data (4.1);
- definition of dependability indicators (A.2.4.3).

In [5] (GOST R 51901-12-2007):

- definition of failure mode (5.2.9);
- definition of failure rate (5.3.4).

In [7] (GOST R 51901-14-2007):

- definition of probability of no-failure (6.1).

In [9] (GOST R IEC 61165-2019):

- definition of dependability indicators (9.1);
- definition of probabilities of states (9.1);
- defining system characteristics (8.1).
- definition of expressions for the probability of no-failure and time to failure (C.3.2).

“Estimate” and “definition” are simultaneously used in the following cases:

In [4] (GOST R 51901-5-2005):

- definition of numerical estimates of dependability indicators (4.1);
- definition of numerical estimates of reliability indicators (A.1.12.1);
- definition of estimates (A.1.12.1);
- defining the estimates of dependability indicators (A.1.12.1).

In [7] (GOST R 51901-14-2007):

- definition of numerical estimates of reliability indicators (6.1).

In [9] (GOST R IEC 61165-2019):

- definition of availability estimates (B.1);
- definition of dependability indicators (Section 1);
- definition of probability of failure (7.2).

Let us clarify the meanings of the terms “definition” and “estimate”. In the Standardization Recommendations [12], it is stated that definition is a logical technique that allows distinguishing, finding and representing a relevant concept. This logical technique is a wording that clarifies the meaning, content, essence, primary characteristic features of a term using known and meaningful words. According to [12], a definition is the starting point for selecting an appropriate term as standardized.

It should be noted that, according to standard [2], “estimate” is a statistic used for the purpose of estimating a parameter that is a feature of a family of distributions. “Estimation” is a procedure that helps obtaining a statistical representation of a general population from a random sample obtained from such general population.

In standard [3], those terms are defined similarly. An estimate is a statistic used for estimating a parameter, while an estimation is an operation involving the definition of the numerical values of distribution parameters based on sample data. A statistic is defined as a function of sample values. It is a random value that may take different values from sample to sample. A parameter is defined as a value used in the description of the probability distribution of a random value.

Thus, in accordance with standards [2] and [3], the terms “estimate” and “estimation” are directly associated with mathematical statistics.

Table 1. Terms substantiated in [10]

Terms used in standards	Terms proposed in [10]
1. “Dependability estimate”	“Dependability calculation”
2. Method of dependability estimation; Method of dependability definition	Method of dependability calculation
3. Computational method of dependability definition	Probabilistic method of dependability calculation
4. Experimental method of dependability definition	Statistical method of dependability calculation

Methods

The Standardization Recommendations [12] state that the terminology is to be unambiguous and self-consistent. The Recommendations set forth the requirements a used term is to comply with. A term must express only one concept and a single concept must be expressed by only one term. Two or more definitions of a single concept are not acceptable. Polysemy (homonymy) and synonymy violate this principle.

It is obvious that the terms, definitions and basic concepts used in standards are to comply with the above requirements. Thus, the integrity of terms, definitions and basic concepts, in the author’s opinion, comes down to the following:

- 1) all terms, definitions and basic concept set forth in the standard are to be unambiguous and self-consistent;
- 2) all terms, definitions and basic concepts set forth in the standard are to be consistent with other national standards and not contradict preceding standards.

That means that a national standard **must not contain**:

- 1) different terms, definitions and basic concepts with identical scope and meaning (must not contain synonyms);
- 2) one and the same term, definition and basic concept with different scope and meaning (must not contain homonyms);

3) terms previously adopted in other national standards with new, modified scope.

When the above deviations and discrepancies are the case, the standard must contain the required explanations and justifications.

In accordance with standards [2] and [3], an estimate of a dependability indicator is a numerical value of a dependability indicator calculated from sample data, while an estimation of a dependability indicator is an operation of obtaining (calculating) the numerical values of a dependability indicator from sample data. Thus, an estimate of a dependability indicator (or another parameter) is a random variable that can take different values from sample to sample. An estimation is done using statistical methods with the aim of obtaining an estimate of a dependability indicator (or parameter).

In this context, the terms “reliability estimate”, “tree estimate”, “gate estimate”, “dependability estimate” refer to properties, diagrams, devices and have special meanings that should be clarified.

In standard [4], Annex A.1.12, section “Statistical methods of estimating the probability of no-failure”, the terms “estimate of probability of no-failure”, “estimate of reli-

Table 2. Modification of the terms defined through “estimate”

Term used in the standard	Proposed term
1. Estimate of probability of no-failure	Calculation (computation) of probability of no-failure
2. Estimate of dependability indicators	Calculation (computation) of dependability indicators
3. Dependability indicators are estimated based on probabilities	Dependability indicators are calculated based on probabilities
4. Estimate of performance and maintainability indicators	Calculation (computation) of performance and maintainability indicators
5. Evaluated indicators	Computed indicators
6. Estimate of probabilistic characteristics	Calculation (computation) of probabilistic characteristics
7. Estimates based on state-transition diagrams	Calculations based on state-transition diagrams (using a state-transition diagram)

Table 3. Modification of terms defined through “definition”

Term used in the standard	Proposed term
1. Definition of availability coefficient	Calculation (computation) of availability coefficient
2. Definition of method	Substantiation (selection) of method
3. Defining average duration of states	Calculating (computing) average duration of states
4. Defining frequency of states	Calculating (computing) the frequency of states
5. Expression for defining average operation time	Expression (formula) for calculating average operation time
6. Formulas for defining dependability indicators	Formulas for calculating (computing) dependability indicators
7. Formulas for defining failure rate	Formulas for calculating (computing) failure rate
8. Probability definition	Calculation (computation) probabilities

ability indicators”, “estimate of dependability indicator” are used correctly. However, at the same time, the use of the term “estimate” has nothing to do with statistics: “estimate of dependability indicators” (4.1), “estimate of dependability improvement” (4.1), “estimate of event characteristics” (4.4.1).

In most standards, in the “Terms and definitions” section, the concept of “definition” is used correctly, i.e., terms are defined.

However, in other cases, the concept of “definition” is used in a sense different from [12], i.e., nothing is actually defined.

It should be noted that in some standards the terms “estimate” and “definition” are used together. A simultaneous use of these terms is difficult to understand.

The meaning and content of some terms were examined in [10]. In this paper, the author relies on the terms substantiated in [10] and set forth in Table 1.

Results.

Thus, based on the requirements for term integrity and the terms examined in [10], the author proposed wordings of terms instead of those using the words “estimate” and “definition”. Those wordings are shown in Tables 2 and 3.

The definitions of the terms in Table 1 are given in [10]. The meaning of terms in Tables 2 and 3 is so obvious that they need no definitions.

Discussion and conclusions

The use of the terms “estimate” and “definition” in some standards lacks integrity. The requirements of the standardization recommendations are not observed, the terms are not unambiguous and consistent with other standards.

It should be noted that along the state standards, these terms are used in technical documentation, science papers, monographs and textbooks.

The options proposed herewith should be considered as a tentative proposal. Final terms are to be developed through extensive discussion and a compromise solution.

References

1. GOST R 27.302-2009. Dependability in technics. Fault tree analysis. Moscow: Standartinform; 2012. (in Russ.).
2. GOST R ISO 3534-1-2019. Statistics – Vocabulary and symbols – Part 1: General statistical terms and terms

used in probability. Moscow: Standartinform; 2020. (in Russ.)

3. GOST R 50779.10-2000. Statistical methods. Probability and general statistical terms. Terms and definitions. Moscow: Standartinform; 2005. (in Russ.)

4. GOST R 51901.5-2005. Risk management. Guide for application of analysis techniques for dependability. Moscow: Standartinform; 2005. (in Russ.)

5. GOST R 51901.14-2007. Risk management. Procedure for failure mode and effects analysis. Moscow: Standartinform; 2008. (in Russ.)

6. GOST R 51901.14-2005. Risk management. Fault tree analysis. Moscow: Standartinform; 2005. (in Russ.)

7. GOST R 51901.14-2007. Risk management. Reliability block diagram and Boolean methods. Moscow: Standartinform; 2008. (in Russ.)

8. GOST R 51901.15-2005. Risk management. Application of Markov techniques. Moscow: Standartinform; 2005. (in Russ.)

9. GOST R IEC 61165-2019. Dependability in technics. Application of Markov techniques. Moscow: Standartinform; 2019. (in Russ.)

10. Zelentsov B.P. Suggestions for improved dependability-related terminology. *Dependability* 2021;21(2): 28-30.

11. Pokhabov Yu.P. On the definition of the term “dependability”. *Dependability* 2017;17(1):4-10.

12. Standardization recommendations R 50.1.075–2011. [Development of standards on terms and definitions]. Moscow: Standartinform; 2012. (in Russ.)

About the author

Boris P. Zelentsov, Doctor of Engineering, Professor of the Department of Further Mathematics, Siberian State University of Telecommunications and Information Sciences, Novosibirsk, Russian Federation, e-mail: zelentsovb@mail.ru

The author’s contribution

The author conducted a terminological analysis of the standards that define the application of the terms “estimation” and “definition”. As a result, terms have been identified that needed clarification.

Conflict of interests

The author declares the absence of a conflict of interests.

Steganalysis of the methods of concealing information in graphic containers

Yaroslav L. Grachev^{1*}, Valentina G. Sidorenko^{1, 2}

¹ Russian University of Transport, Moscow, Russian Federation, ² HSE University, Moscow, Russian Federation
*yaroslav446@mail.ru



Yaroslav L. Grachev



Valentina G. Sidorenko

Abstract. Aim. Today, there is a pressing matter of protection against steganography-based attacks against information systems. These attacks present a danger as they use the most common data files – especially graphics files – as containers that deliver malicious code to a system or cause a leak of sensitive information. Developing methods of detecting such hidden information is the responsibility of a special subsection of steganography, the steganalysis. Such methods should be extensively used in computer forensics as part of security incident investigation, as well as in automated security systems with integrated modules for analysing data files for malicious or dangerous information. An important feature of such activities is the need to examine a wide variety of elements and containing files. In particular, it is required to verify not only the colour values of the pixels in images, but their frequency characteristics as well. This raises a number of important questions associated with the best practices of applying steganalysis algorithms and making correct conclusions based on the outputs. The paper aims to briefly analyse the most important and relevant methods of steganalysis, both spatial and frequency, as well as to make conclusions regarding their performance and ways to analyse the outputs based on the test results of the software that implements such methods. **Methods.** The steganalysis of concealment within the least significant bits of an image's pixels uses Pearson's Chi-square statistical analysis, as well as the Regular-Singular method that involves signature analysis of pixel groups and analytical geometry tools for estimating the relative volume of the hidden message. The Koch-Zhao method of steganalysis is used for the purpose of detecting information embedded in the frequency-domain image representation. It also allows identifying the parameters required for extracting the hidden message. **Results.** A software suite was created that includes the software implementations of the analysed methods. The suite was submitted to a number of tests in order to evaluate the outputs of the examined methods. For the purpose of testing, a sample of images of various formats was compiled, in which information was embedded using a number of methods. Based on the results of the sample file analysis, conclusions were made regarding the efficiency of the analysed methods and interpretation of the outputs. **Conclusion.** Based on the test results, conclusions were made on the accuracy of the steganalysis methods in cases of varied size of the embedded message and methods of its concealment. The patterns identified with the help of the analysis outputs allowed defining a number of rules for translating the outputs into conclusions on the identification of the fact of detection of hidden information and estimation of its size.

Keywords: steganalysis, chi-square method, RS method, Koch-Zhao method, stegocontainer.

For citation: Grachev Ya.L., Sidorenko V.G. Steganalysis of the methods of concealing information in graphic containers. *Dependability* 2021;3: 39-46. <https://doi.org/10.21683/1729-2646-2021-21-3-39-46>

Received on: 12.02.2021 / **Upon revision:** 16.07.2021 / **For printing:** 17.09.2021.

Introduction

Currently, graphics files and data account for a significant share of the network traffic and can be found everywhere, not only as images posted at various network resources, but also as elements of graphical interfaces and design solutions, components of more complex data formats, and many more.

At the same time, the recent years saw a significant growth of malicious software that attacks information systems using steganography, i.e., methods of concealing the fact of transmission of a secret message [1]. Normally, such attacks use image files as containers for delivering potentially hazardous data. That is due to the fact that significant amounts of data can be concealed within them without making any distortions visible to the human eye.

Steganalysis is one of the disciplines of steganography that studies the ways of detecting secretly transmitted information within an analysed information object. Secretly transmitted information is usually understood as information concealed using certain steganographic methods. Additionally, steganalysis also studies the ways and feasibility of extracting concealed information in the process of its detection in the absence of the required input data [2].

Despite the great variety of algorithms of concealing the fact of information transmission in graphics files, almost all of them come down to a number of basic steganographic methods. Those include the method of concealing within the least significant bits of pixels, as well as the Koch-Zhao method that encodes information within an image's representation in the frequency domain [2, 3]. Most other steganographic methods are modifications or variants of those two.

In order to enable the detection of information concealed using the above methods, a number of steganalysis techniques have been developed, whose software implementation allows automating the process of analysis and conducting it without any human involvement (unlike, for example, in the case of various visual attacks).

1. Methods of steganalysis

1.1. The Chi-square method

A method of attacking a stegosystem using Chi-square analysis was proposed and described by Andreas Westfeld and Andreas Pfitzmann in 1999. This method is designed for detecting information concealed through the method of least significant bits (LSB).

First, the concept of pairs of values (PoV) is introduced. Each pair of values is a pair of bytes that encode the colour intensities that differ by only one least significant bit. Essentially, the LSB method performs transformations within such pairs, changing, if necessary, the byte value from the original to the "adjacent" one in the respective PoV [4].

The idea of this method of analysis is based on the assumption that, within an empty container, the probability of a simultaneous appearance of both values of each pair

is low. Therefore, significantly different is the number of colour intensity values that differ by the least significant bit [4]. In other words, for an empty container, the difference in the number of occurrences of the two values of a single pair is significant. Therefore, it is for all PoVs. The number of occurrences of each value is also called frequency.

As the theoretically expected distribution, the sequence of the arithmetic mean frequency values of all pairs is chosen. Since, in case of concealment in the LSB, the frequencies are only redistributed within a pair and the sum of the pair's frequencies remains unchanged. Therefore, the arithmetic mean frequency value within the PoV remains constant.

The observed sample is understood as a sequence consisting of only the even or only odd values of all PoVs, which is due to the requirement of further comparison of the distribution of such samples with the theoretically predicted distribution as part of Chi-square criterion calculation. In this context, the only relevant factor is the difference between a pair's mean frequency and any frequency observed within such pair.

Thus, the theoretically expected sequence of values made up of pair averages is such for both an empty and a populated container. The degree of similarity between the distribution of the observed sample and the theoretically expected distribution thus becomes the measure of the probability of a steganographic embedding within a container. If a Chi-square estimate allows concluding that the deviations from the theoretically expected distribution are insignificant, that strongly suggests the presence of embedded information.

This method of steganalysis is more efficient if applied not to an entire image, but parts of it. In most cases, the image is divided into blocks of about 1% of the total image area or into conventional lines of the pixel matrix. The latter method allows seeing the approximate beginning and end of the sequentially embedded message.

Although the smallest visible distortions are caused by changes of the blue colour channel pixels, the methods of concealing in the LSB allow using all three channels simultaneously due to the fact that the human eye poorly detects colour changes in case of inversion of the least significant bits of a pixel [5]. In this context, the average probability of concealed information should be calculated for all three colour channels of an analysed image block. Even if concealment was only done in a single channel, the average probability will be noticeably non-vanishing and will allow concluding on the presence of concealed information within such block of pixels.

1.2. The RS method

The Regular-Singular method for identifying steganographically concealed messages was proposed by Andreas Pfitzmann, Jessica Fridrich and Miroslav Goljan in 2001. The method is based on the analysis of disjoint groups of n adjacent pixels. n is even [6]. Once the groups have been identified, a regularity function is introduced. That is a function that corresponds a single real number to a single group

and shows the regularity of the group's pixels. The value of the regularity function should be greater the noisier the pixel group is.

As the regularity function, the sum of absolute differences (sum of value differences) of adjacent pixels of a group is chosen:

$$f(G) = f(g_1, g_2, \dots, g_n) = \sum_{i=1}^{n-1} |g_{i+1} - g_i| \quad (1)$$

where G is a pixel group;

g_i is the i -th element of the pixel group G ;

n is the number of pixels in the group.

After calculating the regularity values for all groups of the analysed image, a group of flipping functions is defined. Those functions correspond to the following set of properties:

- 1) $\forall x \in P : F(F(x)) = x, P = \{0, 255\};$
- 2) $F_{dr} : 0 \leftrightarrow 1, 2 \leftrightarrow 3, \dots, 254 \leftrightarrow 255;$
- 3) $F_{inv} : -1 \leftrightarrow 0, 1 \leftrightarrow 2, \dots, 255 \leftrightarrow 256;$

The group of flipping functions F consists of the direct F_{dr} , the inverse F_{inv} and the zero F_0 .

The flipping functions emulate the addition of reversible noise, amplify outliers in a group and reduce its regularity [2].

In order to apply those functions (also called "flipping") to the values of a group's pixels, a mask is used that describes the group of flipping functions applied to the pixel group. A mask is a group of n values, each of which is selected out of three: $-1, 0$ or 1 . Each of them encodes one of the three flipping functions: value " -1 " corresponds to F_{inv} , " 0 " corresponds to F_0 , " 1 " corresponds to F_{dr} . Thus, in the process of flipping, a pixel of a group is subject to a flipping function that corresponds to it in the mask.

Upon the application of flipping functions to a group, the current values of the regularity function are compared to those before the flipping. Based on that comparison, the group belongs to one of the classes: regular, singular, unusable:

- if $f(F(G)) > f(G)$, then $f \in R$ (the group is regular);
- if $f(F(G)) < f(G)$, then $f \in S$ (the group is singular);
- if $f(F(G)) = f(G)$, then $f \in U$ (the group is unusable).

For each group, flipping is done twice, i.e., with a direct and an inverted mask. Upon the classification for all groups, a number of quantitative characteristics are calculated:

- number of regular groups for mask M : R_M ;
- number of singular groups for mask M : S_M ;
- number of regular groups for inverse mask $-M$: R_{-M} ;
- number of singular groups for inverse mask $-M$: S_{-M} .

All the above characteristics are defined as relative values, i.e., as percentages of the total number of groups k . Thus, $R_M + S_M \leq 1$ and $R_{-M} + S_{-M} \leq 1$. The fundamental hypothesis of this method is the assumption that in an empty container, the numbers of single-class groups for the regular and inverse mask are: $R_M \cong R_{-M}$ and $S_M \cong S_{-M}$.

Fig. 1 shows a typical representation of what is called an RS diagram, a graph of values R_M, S_M, R_{-M} and S_{-M} depending on the number of pixels with inverted LSBs in the image [6].

P in Fig. 1 and further refers to the percentage of population of a stegocontainer with a concealed message (relative length of message). Plotted on the x axis is the percentage of pixels with inverted LSBs, plotted on the y axis is the percentage of groups of regular and singular classes of the direct and inverted masks (out of the total number of groups).

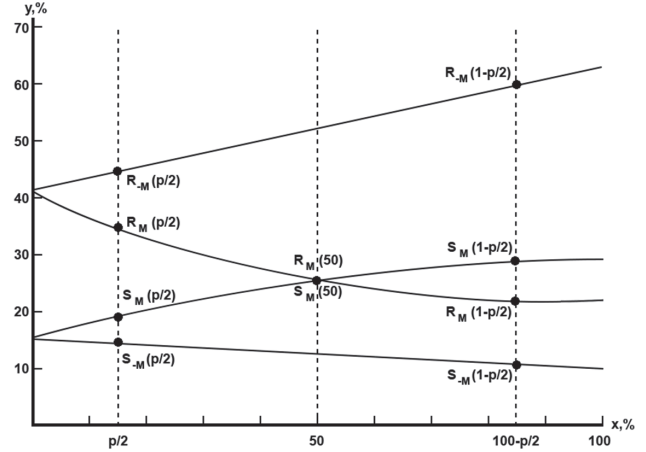


Fig. 1. RS diagram

Based on the information on this typical behaviour of the graphs of quantitative characteristic of the groups, further calculations are performed that allow estimating the relative length of the concealed message.

If the relative length of the concealed message is p , then, since p is a random bit stream, on average, $p/2$ LSBs are inverted in an image. In such case, a 100% population of a stegocontainer causes a situation when $p/2=50\%$. The inversion of a half of the LSBs means that the difference between the number of regular and singular groups will come down to zero. Then $R_M \cong S_M$, which can be seen in the RS diagram.

The numerical measurements of the groups correspond to the points of the RS graph $R_M(p/2)$, $S_M(p/2)$, $R_{-M}(p/2)$ and $S_{-M}(p/2)$. The calculated numerical characteristics of the groups for the same image after all of its LSBs have been inverted will correspond to points $R_M(1-p/2)$, $S_M(1-p/2)$, $R_{-M}(1-p/2)$ and $S_{-M}(1-p/2)$.

Then, the RS method suggests approximating the curves that pass through points $R_{-M}(p/2)$, $R_M(1-p/2)$ and $S_{-M}(p/2)$, $S_M(1-p/2)$ respectively, with straight lines. The curves that pass through points $S_M(p/2)$, $S_M(1-p/2)$ and $R_M(p/2)$, $R_M(1-p/2)$ are approximated with square parabolas subject to the existing points of intersection of lines and parabolas. A parabola and a line that correspond to the same mask have an intersection point on the x axis of the RS diagram, while the parabolas, as it was mentioned above, intersect if $p/2=50\%$.

In order to estimate the relative length of the message p , a system of 11 equations with 11 unknown variables must be solved. Such variables are two coefficients for each of the two lines, three coefficients for each of the two parabolas and the value p . The system includes 8 straight-line or parabola equations for the 8 previously found points, 2 equations for the intersection points of parabolas and corresponding straight lines and an equation for the intersection

point of parabolas. The system's solution allows finding the value p [6].

The key feature of the RS method is that it analyses the quantitative characteristics of small groups of pixels. Due to that, it, while not being able to detect the area of potential embedding, can detect a concealment made in random bits, rather than sequentially.

1.3. Steganalysis using the Koch-Zhao method

This type of analysis is intended for detecting messages embedded into container images using the Koch-Zhao method. The method searches for information encoded in the frequency-domain representations of images.

The frequency-domain representations of an image are generated by calculating the discrete cosine transform (DCT) coefficients, for which the image is divided into blocks of 8×8 pixels, upon which a bivariate DCT is performed on each block producing a matrix of 64 coefficients [7].

In the resulting matrix, the coefficient in the upper left corner that corresponds to the zero frequency (matrix element indexed $(0; 0)$) is called the *DC* coefficient. It defines the primary shade (average colour intensity) of the entire block. All other resulting coefficients are called *AC* coefficients and express the frequency of colour intensity variation along different directions in the selected block (horizontal and vertical) [8].

Thus, each matrix of DCT coefficients is divided into three subsets, i.e., low-frequency, mid-frequency and high-frequency (from the upper left to the lower right corner of the matrix).

The low-frequency coefficients have a greater effect on the colour intensity of pixels. Accordingly, any variations and transformations to the DCT coefficients are done in the mid- or high-frequency regions.

One of the most important tasks that needs to be addressed when attempting the detection of a Koch-Zhao embedding is to make a correct conclusion regarding which DCT coefficients were used for such embedding. Since the application

of the Koch-Zhao method involves concealing information in one of the sets of mid-frequency components, the basic analysis operations are performed for each of these sets individually [9].

Since the bits of the concealed message are encoded through the difference between the absolute values of the selected coefficients, the absolute values of such differences for all the image blocks must be calculated first:

$$C_i = \left| |D_i(k1, k2)| - |D_i(k2, k1)| \right|, \quad (2)$$

where $D_i(x, y)$ is the value of the DCT coefficient indexed (x, y) in the i -th block.

What is calculated at this stage is not the difference between the absolute coefficient values, but rather their absolute values. That is due to the fact that bit encoding is defined by crossing the threshold value P for zero and $-P$ for one [9].

Despite the possibility of fluctuations in the form of different peak values C_i in blocks that were not used to encode bits of a concealed message, the actual blocks used for embedding are recognized by the relatively long continuous section of peak values.

Fig. 2 shows the histograms of values C_i for the same empty and populated container, respectively. Along the x axis are the indexes of blocks, along the y axis are the values C_i of the corresponding blocks.

As we can see, the histograms of the dependence of value C_i on the block number for an empty and populated containers will differ in that the latter has a step-like section.

For the purpose of obtaining the limits of the area containing the embedding, a sequence of modules of the histogram value differences is made, where the two largest values should correspond to the limits of the area of embedding.

However, if the image is large, detecting a small concealed message (especially one encoded with a low threshold value) may be significantly complicated due to the "random" peak values of sequence C_i . That may be due to the use of graphics formats that allow compression, the differences in the accuracy of DCT calculation on the encoding and

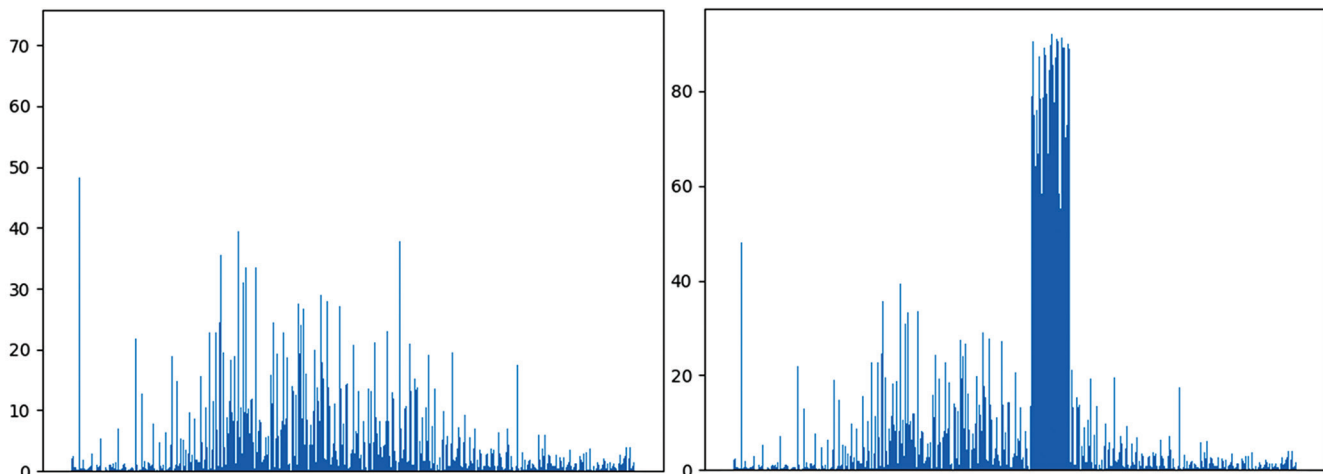


Fig. 2. Histograms of values C_i of an empty and populated container

analysing devices, as well as possible distortions and losses in the process of image transmission.

In order to reduce the probability of failure of the algorithm of embedding detection due to such “noise”, the histograms are to be pre-analysed taking into account the fact that the embedded message corresponds to the continuous section of the peak values, i.e., to find the longest such section.

Thus, it can be concluded regarding the presence of an embedded message and the limits of the concealed message can be defined from the block index values. In order to calculate the encoding threshold, the minimal value out of C_i must be found within the detected area of embedding. This value can be used to retrieve information.

The above analysis procedures are to be performed for each of the assumed pairs of DCT coefficients. Normally, those are the following sets of coefficient pairs: (3;4) and (4;3), (3;5) and (5;3), as well as (4;5) and (5;4) [8]. Out of the obtained results, the one should be selected that corresponds to the highest detected value of the encoding threshold P .

2. Results of method performance testing

The above methods of steganalysis were tested using the software suite developed by the authors that implements all three detection algorithms.

As part of the testing program, a sample of 1600 image files was analysed, in which, using one of the three methods (serial variation of the LSB method, pseudo-random variation of the LSB method, the Koch-Zhao method), one of nine earlier prepared text files was concealed or no information was concealed at all.

For each sample file, an analysis was performed using all three implemented methods: Chi-square, RS and steganalysis of Koch-Zhao concealment. The methods’ performance in the form of relative (for the first two) or the absolute (for the last) estimated length of the message concealed in the image were compared with the actual sizes of the message concealed in each container.

Average deviation means the arithmetic average of the absolute values of the difference between the detected relative size of the concealed image and the actual relative size. The mean deviations for the Chi-square and RS methods are shown in Table 1.

Table 1. Average deviations for the Chi-square and RS methods

Name of method	For empty containers, %	Sequential LSB, %	Pseudorandom LSB, %
Chi-square	0.105	1.411	22.225
RS	4.055	22.996	4.403

As it can be seen, the Chi-square method has almost no false positives on empty containers, while RS, in average, shows an extremely low level of occupancy, about 4%.

It can also be seen that the Chi-square method demonstrates good accuracy when detecting sequential embedding in the least significant bits (on average, the error of estimation of the message size is only about 1.4%), while in case of pseudorandom embedding the error is slightly less than a quarter of the size of the container, i.e., only allows detecting or suspecting the fact of embedding.

Regular-Singular, on the contrary, demonstrates good accuracy in identifying pseudorandom concealment and extremely poorly estimates the size of sequentially embedded messages.

In order to better evaluate the efficiency of hidden message size estimation using Chi-square and RS, the mean deviations for various rates of relative container occupancy were calculated. Those values are shown in Tables 2 and 3.

Table 2. Mean deviations for various occupancy rates for the Chi-square method

Occupancy rate	Mean deviation in case of sequential LSB, %	Mean deviation in case of pseudorandom LSB, %
100%	0.609	0.000
90% – 100%	2.107	26.665
80% – 90%	3.678	61.536
70% – 80%	2.851	69.184
60% – 70%	5.013	54.700
50% – 60%	3.248	52.638
40% – 50%	1.783	44.272
30% – 40%	1.761	35.669
20% – 30%	1.635	24.164
10% – 20%	1.019	13.543
0% – 10%	0.366	5.384

Based on the above findings, a number of patterns can be identified. The accuracy of identification of the size of a message that is concealed using the sequential method is quite high for Chi-square in case of any message size. The largest deviation of about 5% is observed in cases when containers are 60-70% full.

At the same time, Chi-square enables an extremely accurate identification of the size of a concealed pseudo-random message if it is 100% of the size of the container. If it is below 100%, but above 90%, the method definitely identifies an embedding, but the error of message size estimation is significant. If the occupancy is below 90%, Chi-square is not always able to detect a concealment. Its average deviation in some cases is roughly equivalent to the relative message size. If the occupancy is below 60%, Chi-square shows the worst results in case of pseudo-random concealment. If the occupancy is between 60% and 90%, it is able to identify embedding, but estimates the message size with enormous errors.

Table 3. Mean deviations for various occupancy rates for the Chi-square method

Occupancy rate	Mean deviation in case of sequential LSB, %	Mean deviation in case of pseudorandom LSB, %
100%	70.501	11.593
90% – 100%	67.184	8.060
80% – 90%	57.385	4.112
70% – 80%	50.704	1.806
60% – 70%	40.795	1.638
50% – 60%	29.291	1.264
40% – 50%	21.286	2.977
30% – 40%	11.476	2.985
20% – 30%	6.277	2.166
10% – 20%	5.254	3.832
0% – 10%	3.894	3.722

RS shows good accuracy of pseudo-random embedding identification. However, if the occupancy is below 10%, it is not always able to identify a concealment (the average deviation is about 3.7%, while the deviation for empty containers is about 4.4%). If the occupancy is between 20% and 80%, the method demonstrates high accuracy (less than 3% of average deviation). If the size of a pseudo-randomly concealed message is above 80%, the error of the RS method is noticeably higher. It is the greater the larger is the size of the embedded information. The method demonstrates a higher error when the occupancy is 10-20%.

While evaluating the size of a sequentially concealed message, RS shows far worse results and has an average deviation the higher the higher the occupancy rate. At the same time, a message size below 10% in most cases prevents a clear identification of a fact of concealment (same as with pseudo-randomly concealed information and taking into account the average deviation for empty containers).

If the occupancy is above 10%, RS will estimate the size of the hidden message on average at about 15-30% of the container size, which allows making a clear conclusion of the fact of embedding, yet its size is identified very inaccurately.

In order to evaluate the performance of the algorithm of steganalysis of Koch-Zhao concealments, the number of correctly recognized images was identified, for which the algorithm correctly estimated the size of the hidden message, the incorrectly recognized, for which the algorithm retrieved an incorrect, yet non-zero message size value, and the unrecognised, for which the algorithm retrieved 0 in a situation when the image contained actual concealed information.

The above characteristics are presented in Table 4.

Using the Koch-Zhao method, all the empty containers were correctly identified as images that do not contain concealed information.

As it can be seen from the table, in more than 90% of cases the algorithm correctly identified the size of the embedded message, and therefore, it could be likely extracted. In most other cases, the size was identified incorrectly, which prevents the extraction of concealed data (at least without manual operations), yet it allows making a clear conclusion regarding the presence of concealment in the frequency domain. 1.3% of the images remained completely unidentified, which may be evidence of a low chosen information embedding threshold or high level of noise that causes outlying DCT matrix value differences that complicate the identification of the embedding area.

The analysis of the test results also showed that the Koch-Zhao steganalysis did not reveal a single case of concealment in the spatial domain of an image (in the least significant bits), while the Chi-Square and RS failed to produce a result that would allow identifying the fact of embedding in the image representation in the frequency domain.

3. Conclusions regarding the use of the methods of steganalysis

The above findings show that in order to obtain clear conclusions regarding the presence of embedded information into least significant pixels of an image, both Chi-Square and RS must be used. A comparison of the outputs of both methods in most cases allows identifying the type of embedding, i.e., sequential or pseudo-random,

Table 4. Characteristics of various result types of the Koch-Zhao concealment analysis

Result type	Number	Relative number
Correctly identified	277.	92.3%
Incorrectly identified	19.	6.3%
Not identified	4.	1.3%

as well as more or less accurately estimating the size of the message.

If the Chi-square estimate is below 0.1% and the RS estimate is below 4%, there is no embedding and the container is empty.

If the Chi-square estimate is close to 100%, the container is fully or almost populated. If the RS estimate is about 30% or less, the information is embedded sequentially, if it is about 80% or more, it is embedded pseudo-randomly.

An RS estimate of more than 30% and Chi-square estimate not higher than 30% indicate a pseudo-random embedding. If the RS estimate is less than 80%, it can be considered an almost accurate estimate of the size of the concealed message. Otherwise, it can be reliably assumed that the size of the concealed information exceeds 80%.

If the RS estimate is below 30%, it should be compared with the Chi-square estimate. The latter being times lower or close to zero indicates a pseudo-random embedding with the message size corresponding to the RS estimate. However, if the Chi-square estimate is approximately equal to or greater than the RS estimate, a sequential concealment of information is identified. Accordingly, the size of the concealed data should be considered equal to the Chi-square estimate.

The above rules of output comparison and analysis used in the shown sequence allow covering most combinations of the Chi-square and RS outputs. Outputs that do not fit the rules indicate fluctuations that may be caused by a variety of factors (high image noise, extremely small image size, etc.) and require manual intervention to produce the final output.

The obtained conclusions are graphically represented in Fig. 3. p is the final estimate of the size of the concealed message. Shown in blue are the areas of pseudo-random embedding (O.2, O.4, O.5, O.7). Shown in green are the

areas of sequential embedding (O.3, O.6). Shown in gray are the areas of criteria values that require further study. The white area (O.1) indicates an empty container.

The implemented method of frequency domain steganalysis in most cases allows detecting concealments made using the Koch-Zhao method, however, the size of the concealed message can in a number of situations be identified incorrectly. Therefore, in some cases, extracting information concealed in the frequency domain may require a manual intervention in the form of expert analysis of the histograms of DCT coefficient difference.

Thus, these rules and conclusions can be used in integrated security systems or individual image analysis modules, automatically detecting potentially hazardous graphics files that can carry embedded malicious code or sensitive data.

Conclusion

The conducted study of the methods of steganalysis enabled their software implementation as a single suite. The suite's testing on a sample of 1600 image files allowed assessing the overall efficiency of the methods, the average errors in the estimation of the length of embedded messages. The tests revealed a number of patterns, a dependence between the outputs of the methods and the size of the hidden messages. The revealed patterns lead to a number of conclusions regarding the performance of the examined methods of steganalysis. They allow correctly evaluating the obtained estimates of the size of hidden information and making conclusions accordingly.

These patterns and conclusions simplify the work of computer forensics experts who use these steganalysis techniques to analyse images. They can also serve as research criteria when graphics files are analysed for hidden information in automated information security systems or their steganalysis modules designed to prevent information leaks through steganographic channels or steganography-based attacks.

As part of the software suite development, the analysed methods of steganalysis were adapted to the existing graphics formats that use the *TrueColor* colour storage format. Besides the detection of concealment in the LSB, estimation of the area and approximate threshold of information embedding into the representation of an image in the frequency domain, the suite also enables automatic attempts of extracting such information even from significantly noisy images. The conclusions regarding the specificity of the examined steganalysis methods in the form of rules for comparing joint analysis outputs allow fine-tuning the criteria of security systems or traffic analysers for the purpose of preventing information security incidents.

References

1. Varnovsky N.P., Golubev E.A., Logachev O.A. [Modern trends in steganography]. Proceedings of the Conference Mathematics and security of information technologies in MSU, October 28-29, 2004. Moscow: MCCME; 2005.

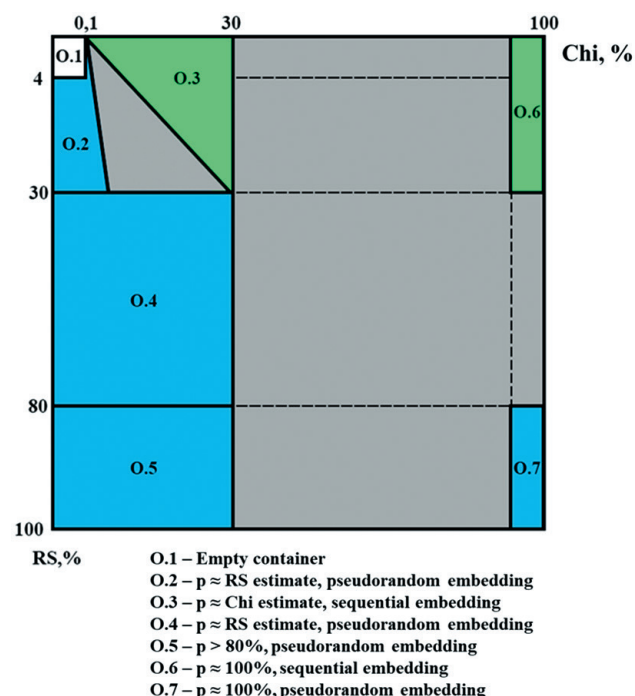


Fig. 3. Graphical representation of outputs

2. Agranovsky A.V., Balakin A.V., Gribunin V.G. et al. [Steganography, digital watermarks and steganalysis: a monograph]. Moscow: Vuzovskaya kniga; 2009. (in Russ.)
3. Konakhovich G.F., Puzyrenko A.Y. [Computer steganography. Theory and practice]. Kiev: MK-Press; 2006. (in Russ.)
4. Westfeld A., Pfitzmann A. Attacks on Steganographic Systems. Dresden University of Technology, Department of Computer Science. Dresden (Germany): 1999. DOI: 10.1007/10719724_5.
5. Gonzalez R.C., Woods R.E. Digital image processing. Moscow: Tekhnosfera; 2005.
6. Fridrich J., Goljan M., Du R. Reliable Detection of LSB Steganography in Color and Grayscale Images. New York: Binghamton University; 2001. DOI: 10.1145/1232454.1232466.
7. Khayam S.A. The Discrete Cosine Transform (DCT): Theory and Application. Michigan: Department of Electrical and Computer Engineering, Michigan State University; 2003.
8. Farid H. Digital Image Forensics. *Scientific American*; 2008.
9. Belim S.V., Vilkhovsky D.E. Koch-Zhao algorithm steganalysis. *Mathematical Structures and Modeling* 2018;4(48):113-119. DOI: 10.25513/2222-8772.2018.4.139-119. (in Russ.)

About the authors

Yaroslav L. Grachev, Student, Russian University of Transport, Moscow, Russian Federation, e-mail: yaroslav446@mail.ru

Valentina G. Sidorenko, Doctor of Engineering, Professor, Chair Professor, Department of Management and Protection of Information, Russian University of Transport, Chair Professor, Department of Business Informatics, HSE University, Moscow, Russian Federation, e-mail: valenfalk@mail.ru.

The authors' contribution

Grachev Ya.L. Development of a software system that implements the analysed methods of steganalysis, sampling, testing and analysis of the results, conclusions on the performance of the tested methods.

Sidorenko V.G. Analysis of methods and principles of steganography, review of the methods of steganalysis.

Conflict of interests

The authors declare the absence of a conflict of interests.

Risk Model of German Corona Warning App – Reloaded

Jens Braband¹, Hendrik Schäbe^{2*}

¹ TU Braunschweig, Braunschweig, Germany

² TÜV Rheinland, Köln, Germany

*schaebe@de.tuv.com



Jens Braband



Hendrik Schäbe

Abstract. Aim. In this paper we discuss the risk model of the German Corona Warning App in two versions. Both are based on a general semi-quantitative risk approach that is not state of the art anymore and for some application domains even deprecated. The main problem is that parameter estimates are often only ordinal scale or rank numbers for which such operations as multiplication or division are not clearly specified. Therefore, it may results in underestimation or overestimation of the associated risk. **Methods.** The risk models that are used in the apps are analyzed. Comparison of the nomenclature of model parameters, their influence on the result, approaches to the generation of a combined risk assessment is carried out. The effectiveness of the models is analyzed. **Results.** It is shown that most of the parameters in the model are used only as binary indicator variable. It has been found that the Corona Warning App uses a much more limited model that does not even assess risk, but relies on one parameter which is weighted exposure time. It has been shown that the application underestimates this parameter and therefore may erroneously reassure users. Thus, it may be concluded that the basic risk model implemented before version 1.7.1., is rather a dosimetric model that depends on the calculated virus concentration and does not depend on exposure and other parameters (excluding some threshold values). It is not even a risk model as defined by many standards. Changes of the risk model in the later version are not fundamental. In particular the later model also assesses not individual risk, but individual exposure according to the results. In addition, the model greatly underestimates the duration of exposure. Although it is reported that about 60% of the app's users have shared positive test results, the absolute number of published results is less than 10% of all positive test results. Therefore, from an individual point of view, the application is effective only in 10% of cases, or even less. **Conclusions.** As the Corona Warning App also has other systematic limitations and shortcomings it is advised not to rely on its results but rather on Covid testing or vaccination. In addition, if there are enough virus tests available in the near future, the application will even become outdated. It will be better to develop an application that can assess risks a priori, as a kind of decision support for its users based on their individual risk profile.

Keywords: Corona Warning App, Risk Model, Model Analysis

For citation: Braband J., Schäbe H. Risk Model of German Corona Warning App – Reloaded // Dependability. 2021. No.3. P. 47-53. <https://doi.org/10.21683/1729-2646-2021-21-3-47-53>

Received on: 24.05.2021 / **Revised on:** 25.07.2021 / **For printing** 17.09.2021

Introduction

In Germany, an app has been introduced to cope with the Covid pandemic: Corona Warning App (CWA). The CWA computes a risk for persons to have been infected caused by contacts. In a former paper we analyzed, how CWA works and also its risk model [12]. Recently, after some criticism, the model has been changed (from version 1.9 and later). We study in this paper the changes of the risk model and its relevance as an indicator for individual risk. While the particular discussion focusses on the German app only, the so-called Exposure Notification Framework (ENF), on which the CWA relies, has been introduced by Apple and Google as part of a standard interface, and so it can be expected that the results may be generalized [14].

We introduce how the CWA works, then briefly summarize the former risk model, describe the changes. Finally, we compare the risk models and assess the effects of the changes.

How the app works

First a simplified overview is given in order to understand the risk model. More details are given in [1] and [2]. A much broader view on contact tracing apps has been published recently [14].

After a user has installed the CWA, each day a new anonymous ID is created. About every five minutes the environment is scanned for Bluetooth signals emitted from other mobiles. Data like ID, signal attenuation, duration etc. are collected and aggregated for each day.

If the user receives a positive test result and agrees to publish it, then the associated anonymous IDs for the preceding 14 days are transmitted to the central server, from which it is transmitted to all subscribers and checked against the recorded data in the local app. The actual risk evaluation is performed decentralized by each CWA.

The Basic Risk Model

The basic model up to version 1.7.1 is defined by four parameters [2], which in a first step are evaluated on a semi-quantitative scale each ranging from 0-8 for each day for each ID that reported a positive test result (see figure 1):

- The Days since Exposure (DE) is the time since exposure to the infected person, a value between 0 and 14, durations longer than 14 days are not taken into account.
- The Exposure Duration (ED) is the cumulative time of exposure on the day, takes values between 0 and 8.

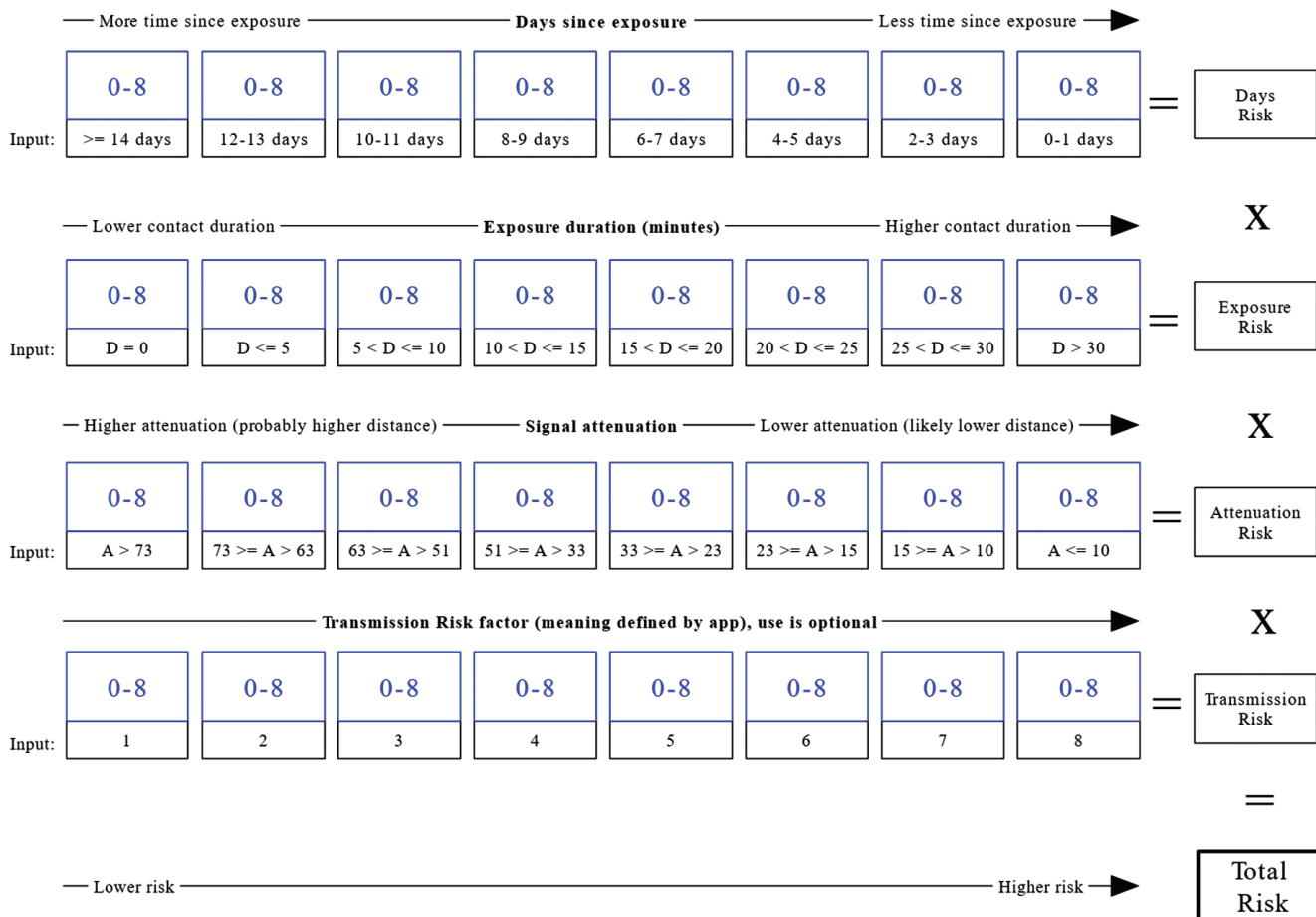


Figure 1: Basic Risk Calculation [2]

• For the 9th DE=11 and so TR=1. However, ED is below the threshold and set to 0. So TRS=0. Otherwise, the TRS would have been 5, which is below the MRS and would have been discarded anyhow.

Combined Risk Model

In a second step each CWA combines the scores for different encounters calculated by the basic risk model. Let $R_1, R_2, \dots, R_n$ denote the individual TRS for different days and different IDs that are above the MRS.

In a first step the maximum value R_{max} of the different TRS is determined. Then the ED of all the n encounters is summed up in three different classes: close, medium and far. Let these durations be t_1, t_2 and t_3 . For each of the classes a weight is defined and additionally a weight offset, which are denoted by w_1, w_2, w_3 and w_4 , respectively. Note that in practice the weights for the close and medium classes outweigh the others. Also, an Average Risk Score (ARS), currently 50 [3], is defined. Then the Total Combined Risk (TCR) is calculated as (see figure 2)

$$TCR = (t_1 w_1 + t_2 w_2 + t_3 w_3 + w_4) \frac{R_{max}}{ARS}. \quad (2)$$

Surprisingly the TCR is in fact not a risk, but an exposure as the result is in minutes. The first term (in brackets) is a weighted exposure time which is adjusted by a relative factor (dimensionless) which depends on the maximum virus concentration compared with some average.

So overall, without full mathematical exactness, we can characterize the approach used until version 1.7.1 [3] by the German Corona warning app as

$$TCR = (t_1 w_1 + t_2 w_2 + t_3 w_3 + w_4) \cdot \frac{\max\{10 \cdot \delta_{ED} \cdot \delta_{DE} \delta_{SA} \cdot TR\}}{ARS} \approx (t_1 + t_2 / 2) \frac{10 \max\{TR\}}{ARS}. \quad (3)$$

Basically, this is not a risk in the narrow sense, it is a weighted exposure duration, the units of TCR are not expected damage per time or similar, but just minutes of exposure. And the weight applied is just a measure of relative infectiosity, which expresses the size of TR relative to a normalizing factor of 5, which is assumed as average infectiosity. And the impact of the factor is limited, the highest possible value being 1.6.

And, if we take additionally into account the uncertainty and spread of all the input parameters, e. g. noise in the signal attenuation, uncertainty in the exposure duration and infectiosity [4], or arbitrariness of the weights chosen, then the model boils down to a quite simple formula and decision procedure:

$$TCR \approx \frac{\max TR}{5} \sum ED, \quad (4)$$

• Estimate the minutes that the person was exposed closely to infected persons ($\sum ED$)

• Weight the exposure ED by the infectiosity of the most infected person ($\max TR/5$)

• Take action if the result is HIGH

Example (continued)

• Additionally, Charlie has received a positive test result on the 20th, but he reported it only on the 21st. But he has installed the app only a week ago.

• He has been on the same bus on the same days, but with some larger distance to Bob (2m)

• For the 16th, DE=5 (evaluated on the 21st) and so TR=5. Both SA and ED are above the threshold and set to 2 and 1, so TRS=50.

• For the 9th, Charlie had not installed the app yet, so there are no data.

• As Alice sat close to Bob, and Charlie in medium distance, $t_1=t_2=20$ minutes. The weights are currently set [3] to $w_1=1, w_2=0.5, w_3=w_4=0$, and so the weighted sum results in 30 minutes

• So the TCR amounts to $30 \times 80 / 50$, which gives 48 minutes.

• The warnings of the app are issued based on the TCR, which are configured [3] as LOW for values up to 15, and HIGH from 15 onwards. So finally Bob would get a HIGH risk warning.

Note that the TCR is almost independent of the distance measured by Bluetooth Signal Attenuation. There is only a loose threshold defined and the exposure duration is weighted into two distance classes. However, it is known, see e. g. the FAQ by RKI [9], that transmission is mainly by aerosols, and the risk increases when the exposure distance decreases. Close contact to infected persons is also known to be a factor in superspreading events. The importance of the distance is also supported by the fact that it is a major factor in the German Corona rules AHA, where the first letter stands for Abstand=distance.

Updated Risk Model

From version 1.9 on the risk model has been changed several times [10], with the main intent of improving the granularity of the risk assessment and also the accumulation of risk i. e. several low-risk encounters in the former model may constitute a high risk in the new model. As a caveat it should be noted that some parameters seem to have been adapted several times and that even on the GitHub repository [11] the parameters in the documentation are not consistent. The information presented here is from end of March 2021 and in doubt the authors have used the parameters directly from the configuration files of the repository rather than from the documentation.

As a main means to achieve this 30-minute evaluation windows are introduced, and a contact is only counted if within this window

1. The signal attenuation (SA) is below the threshold of 73 dB for at least 5 minutes (ED), and

Transmission Risk Level (TRL) to Value (TRV)	
I	0.0
II	0.0
III	0.6
IV	0.8
V	1.0
VI	1.2
VII	1.4
VIII	1.6

Figure 3: Direct mapping of TRL to scaling factor [11]

2. Transmission Risk (TR) level must be at least 3 (compare figure 1)

If these criteria are not met, then the contact is discarded. Compared to formula (1) this seems not a big change but for the second condition, which was not present before.

Close encounters with SA below 55dB will be counted full, while all other encounters will be discounted by 50%. Note that this does not match with the interface defined by figure 1, but the weight used are the same as before, compare formula (2).

For the TR level a new weighting is introduced, ranging from 0.6 to 1.6, see figure 3. This seems to be a shortcut compared to the calculations in formula (2). Most of the factors are the same as before, e. g. for TR=8, 6, 5, 3, but some have been discarded or slightly changed.

Finally, the measured encounter times (more than 5 minutes within a 30-minute window) are weighted by their proximity and multiplied by the TR level weight (from figure 3) and summed up for a particular day. If the sum is below 15 minutes, but above 5 minutes (some sources suggest 13 minutes [11]) the overall risk is LOW or GREEN, otherwise it is HIGH or RED. This is evaluated per day and the app displays the number of days with low or high risk instead of the numbers of encounters with low or high risk (as up to version 1.7.1).

Taking this all into account, it seems as if the general semi-quantitative framework ENF as defined in figure 1 is not really applied anymore. Instead for each encounter i of a particular day a Transmission Risk Score similar to (1) is evaluated (with some changes to the indicator functions):

$$TRS_i = \delta_{ED} \cdot \delta_{SA} \cdot ED \cdot TR, \quad (5)$$

where ED is already understood as the weighted encounter times similar to (2)

$$ED = t_1 w_1 + t_2 w_2 + t_3 w_3 + w_4. \quad (6)$$

Finally, the risk is combined per day by simple summation

$$TCR = \sum_{\text{per day}} TRS_i. \quad (7)$$

Example (reloaded):

- For the encounter with Alice on the 16th DE=4 and so originally TR=8, which by figure 3 is then recoded to TR=1.6. Both SA and ED are above the threshold and the weight is 1 and both times are counted full, respectively, so both encounters score TRS=16 minutes, individually. Summing them up a RED warning of HIGH risk is displayed to Bob for this day.

- For the 9th DE=11 and so TR=1, which is recoded to TR=0. So for this risk TRS=0 even as the other thresholds are fulfilled. Otherwise with the lowest TR=0.6 the TRS would have been 6, which would have been at least a GREEN warning with LOW risk for this day.

- Note that the RED warning is independent from the encounter with Charlie on the 16th. For this day Charlie's parameters are the DE=5 with TR=5 recoded to TR=1. As the distance was larger the ED is halved to 5, so both encounters would score TRS=5, summing up to 10 for the day, which standalone would lead to a GREEN warning and combined with Alice's scores would lead to combined score of 26.

If we take into account that some scaling factors have been left out in comparison with the previous approach, e. g. for DE and SA, then the results in the example are quite comparable. The most obvious differences are that

- the term related to $\max\{TR\}$ in formulae (3) and (4) has been simplified and has been substituted by a fixed factor based on TR as in figure 3, and
- that the risks are accumulated and reported per day and not combined for all encounters as in formula (4)

Discussion

The changes in the risk model are not fundamental, so the general criticism from [12] remains valid. In particular the CWA does not estimate individual risk, but individual exposure, as already the units of the resulting TCR shows, which is [min] and not expected harm or expected severity or similar expressions of risk.

Also, the CWA grossly underestimates the exposure duration for the following reasons:

- There have been about 25 million downloads of CWA so far [13], but it is unknown how many installations are active. If we assume optimistically that 30% of the population would have an active app, then less than 10% of the encounters could be registered, as both parties need to have the app activated. Studies indicate [14] that from 60% coverage of the population by CWA, it could become an effective tool.

- Although it is reported that about 60% of the app users have shared positive test results [13], the absolute number of about 250,000 shared results is again less than 10% of all positive test results (about 2,820,000)

- A large number of short contacts (below 5 minutes), that could be close and infectious, are not detected and registered by the app due to energy saving reasons the app scans only about every 5 minutes its environment

This holds also from an individual user's perspective: a user with active CWA meets only other users with active CWA in less than 30% of the encounters, many short encounters, e. g. in the supermarket or at work, will be discarded and only in 60% of the possible cases other infected users will report their positive results. So, also from the individual perspective the CWA is effective only in about 10% of the cases, or even less.

The Good, the Bad and the Ugly, or: Summary

The risk model of the German CWA has several interesting, somewhat puzzling and also disturbing properties:

1) It has increased the number of users and is also compatible with apps from some neighboring countries, but its coverage is still rather limited.

2) The main advantage is that the CWA records contacts automatically that would otherwise perhaps not have been noticed by the user. But due to data privacy the user can't learn from the warning as it is not reported when and how long the exposure has been. Therefore the CWA can also not support contact tracing.

3) Also the number of positive tests reported is only a small percentage of all positive tests

4) In the narrow definition of many standards, it is not a complete risk model, as it estimates only single parameters of risk, but not a comprehensive risk, more a kind of exposure measure. A partial explanation can be based on the decentralized architecture of the app and the incomplete and inaccurate information it uses

5) Even worse, it does not cover many situations with short exposure that might still go along with a high dose of virus transmission. This may be due to the pressure by the mobile industry to save energy.

6) Also, it is not an a priori risk assessment in order to decide whether a particular action is acceptable, it assesses risk only a posteriori, when it can't be reduced anymore.

7) And finally, most important, it underestimates the true exposure time or virus dose by a large factor, possibly 5 to 10, so it reassures its users to feel safe if it states no or only a low risk, but in reality, the exposure maybe much higher.

In summary, the authors would recommend not to rely on the results of the CWA due to the many shortcomings of its approach and its risk model, some of which are systematic and can not be improved by further updates. A test or even better, vaccination, should always be preferred. Also, it seems that in the near future, when sufficient Corona tests are available, the CWA will even become obsolete. It is rather recommended to develop an app that may estimate risks a

priori as a kind of decision support for its users based on their individual risk profile.

References

1. RKI: So funktioniert die Corona-Warn-App im Detail. Available at: https://www.rki.de/DE/Content/InfAZ/N/Neuartiges_Coronavirus/WarnApp/Funktion_Detail.pdf (accessed at 06.10.2020).
2. Команда CWA: Corona-Warn-App Solution Architecture. Available at: https://github.com/corona-warn-app/cwa-documentation/blob/master/solution_architecture.md (accessed at 06.10.2020).
3. CWA Team: Wie ermittelt die Corona-Warn-App ein erhöhtes Risiko? Available at: <https://github.com/corona-warn-app/cwa-documentation/blob/master/translations/cwa-risk-assessment.de.md> (accessed at 06.10.2020)
4. CWA Team: Эпидемиологическая мотивация уровня риска передачи, 2020-06-15. Available at: https://github.com/corona-warn-app/cwa-documentation/blob/master/transmission_risk.pdf (accessed at 06.10.2020).
5. Bowles J.: An Assessment of RPN Prioritization in a Failure Modes Effects and Criticality Analysis // Proc. RAMS2003, Tampa, January 2003.
6. Braband J. Improving the Risk Priority Number Concept // Journal of System Safety. 2003. No. 3. P. 21–23.
7. Durivage M. Is It Time To Say Goodbye To FMEA Risk Priority Number (RPN) Scores? // Pharmaceutical Online. Available at: <https://www.pharmaceuticalonline.com/doc/is-it-time-to-say-goodbye-to-fmea-risk-priority-number-rpn-scores-00012020-04-27> (accessed at 06.10.2020).
8. Braband J. Beschränktes Risiko // Qualität und Zuverlässigkeit. 2008. No. 53(2). P. 28–33.
9. Robert-Koch-Institut: SARS-Cov-2 Steckbrief. Available at: https://www.rki.de/DE/Content/InfAZ/N/Neuartiges_Coronavirus/Steckbrief.html (accessed at 06.10.2020).
10. SAP Deutschland: Risikoberechnung der Corona-Warn-App nach detaillierten Tests weiter angepasst, 23.02.2021. Available at: <https://www.coronawarn.app/de/blog/2021-02-23-corona-warn-app-risk-calculation-optimization/> (accessed at 06.10.2020).
11. Available at: <https://github.com/corona-warn-app> (accessed at 06.10.2020).
12. Braband J., Schäbe H. Analysis of the Risk Model of German Corona Warning App, Reliability: Theory & Applications. 2021. No. 1 (61). Vol. 16., – P. 62-70
13. Robert-Koch-Institut: Kennzahlen der Corona-Warn-App, 17.02.2021. Available at: https://www.rki.de/DE/Content/InfAZ/N/Neuartiges_Coronavirus/WarnApp/Archiv_Kennzahlen/Kennzahlen_18022021.pdf?__blob=publicationFile (accessed at 30.03.2020).
14. Landau S. People Count – Contact Tracing Apps and Public Health. The MIT Press, 2021.

About the authors

Jens Braband, Dr. rer. nat., Principal Key Expert for RAMSS at Siemens Mobility GmbH, and Honorary Professor, TU Braunschweig, Germany. Email: j.braband@tu-braunschweig.de

Hendrik Schäbe, Dr. rer. nat. habil., Chief Expert on Reliability, Operational Availability, Maintainability and Safety, TÜV Rheinland InterTraffic, Cologne, Germany, e-mail: schaebe@de.tuv.com

The authors' contribution

The contribution of the authors consists of an analysis of the risk models of the Corona applications. The analysis has been carried out by Jens Braband, Hendrik Schäbe has supported with the analysis of the model and some considerations on the application.

Conflict of interests

The authors declare the absence of a conflict of interests.

Intelligent methods for improving the accuracy of prediction of rare hazardous events in railway transportation

Olga B. Pronevich^{1*}, Mikhail V. Zaytsev¹

¹JSC NIIAS, Moscow, Russian Federation

*obpronevich@gmail.com



Olga B. Pronevich



Mikhail V. Zaytsev

Abstract. The paper **Aims** to examine various approaches to the ways of improving the quality of predictions and classification of unbalanced data that allow improving the accuracy of rare event classification. When predicting the onset of rare events using machine learning techniques, researchers face the problem of inconsistency between the quality of trained models and their actual ability to correctly predict the occurrence of a rare event. The paper examines model training under unbalanced initial data. The subject of research is the information on incidents and hazardous events at railway power supply facilities. The problem of unbalanced data is expressed in the noticeable imbalance between the types of observed events, i.e., the numbers of instances. **Methods.** While handling unbalanced data, depending on the nature of the problem at hand, the quality and size of the initial data, various Data Science-based techniques of improving the quality of classification models and prediction are used. Some of those methods are focused on attributes and parameters of classification models. Those include FAST, CFS, fuzzy classifiers, GridSearchCV, etc. Another group of methods is oriented towards generating representative subsets out of initial datasets, i.e., samples. Data sampling techniques allow examining the effect of class proportions on the quality of machine learning. In particular, in this paper, the NearMiss method is considered in detail. **Results.** The problem of class imbalance in respect to the analysis of the number of incidents at railway facilities has existed since 2015. Despite the decreasing share of hazardous events at railway power supply facilities in the three years since 2018, an increase in the number of such events cannot be ruled out. Monthly statistics of hazardous event distribution exhibit no trend for declines and peaks. In this context, the optimal period of observation of the number of incidents and hazardous events is a month. A visualization of the class ratio has shown the absence of a clear boundary between the members of the majority class (incidents) and those of the minority class (hazardous events). The class ratio was studied in two and three dimensions, in actual values and using the method of main components. Such “proximity” of classes is one of the causes of wrong predictions. In this paper, the authors analysed past research of the ways of improving the quality of machine learning based on unbalanced data. The terms that describe the degree of class imbalances have been defined and clarified. The strengths and weaknesses of 50 various methods of handling such data were studied and set forth. Out of the set of methods of handling the numbers of class members as part of the classification (prediction of the occurrence) of rare hazardous events in railway transportation, the NearMiss method was chosen. It allows experimenting with the ratios and methods of selecting class members. As the results of a series of experiments, the accuracy of rare hazardous event classification was improved from 0 to 70-90%.

Key words: machine learning, rare events, class imbalance, better accuracy of predictions, data sampling, class balancing.

For citation: Pronevich O.B., Zaytsev M.V. Intelligent methods for improving the accuracy of prediction of rare hazardous events in railway transportation. *Dependability* 2021;3: 54-64. <https://doi.org/10.21683/1729-2646-2021-21-3-54-64>

Received on: 05.07.2021 / **Revised on:** 02.08.2021 / **For printing:** 17.09.2021.

1. The relevance of the problem of improving the quality of simulation as part of predicting rare events in railway transportation

When predicting hazardous events, researchers face a controversial problem: the fewer are the events, the better it is for the observed object, and the harder it is to train a model of the required quality. Currently, the proportion of hazardous events out of all incidents involving railway power supply facilities (RPSF) does not exceed 2% per year (Fig. 1a).

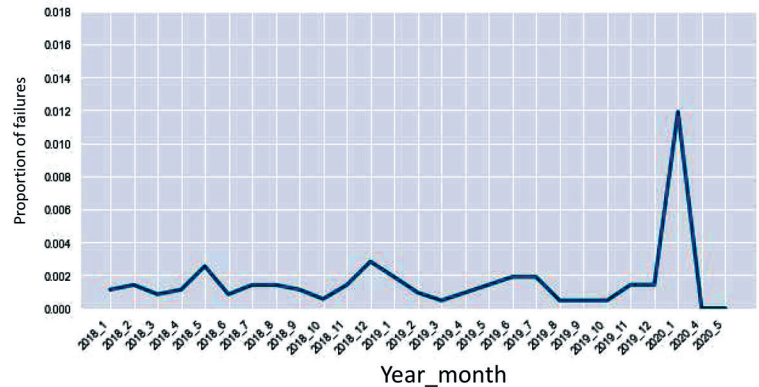
As it can be seen from the graph in Fig. 1a, since 2018, the proportion of hazardous events has been on a decline, yet the available observation period is not long enough to conclude on a steady reduction, while the data up to 2018 suggest the possibility of growing numbers of failures. Detailed monthly statistics (Fig. 1b) show the numbers of hazardous events peaking. The lack of clear seasonal patterns in the data does not allow scheduling preventive measures aimed at prevent-

ing hazardous events. A year is not an efficient observation period, as the condition of railway facilities changes greatly over a year and the planned activities may prove to be irrelevant. Therefore, of special interest are models that allow predicting hazardous events within periods of one month. At this level of observation detail, the matter of the “rarity” of the target factor becomes even more critical. On average, it accounts for less than 0.4% of cases.

Despite the small number of hazardous events, the number of other unrelated incidents is in the thousands, which allows employing Big Data techniques. There are two main obstacles to building highly accurate models, i.e., class imbalance and data quality. In terms of data quality, the most common issues include incompleteness, duplication, inconsistent information, manual input errors. However, when working with railway facilities, the authors came to face another problem, i.e., the lack of distinct differences between the characteristics of facilities. On the monthly level, RPSF are mainly characterized by incident data. Upon data preparation, each RPSF is characterized by more than 150 features.

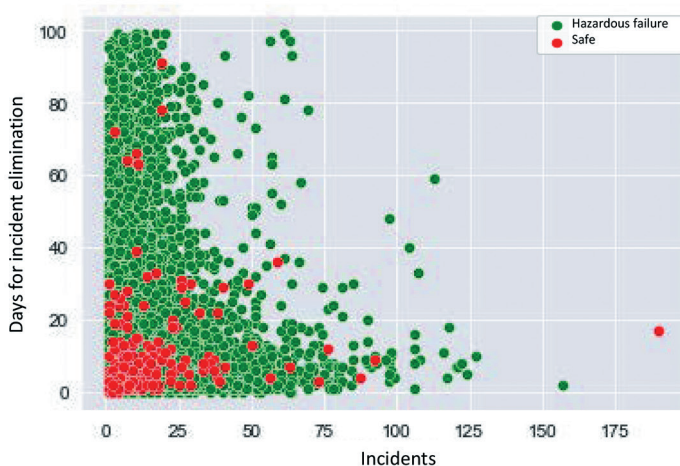


a) proportion of hazardous events year to year

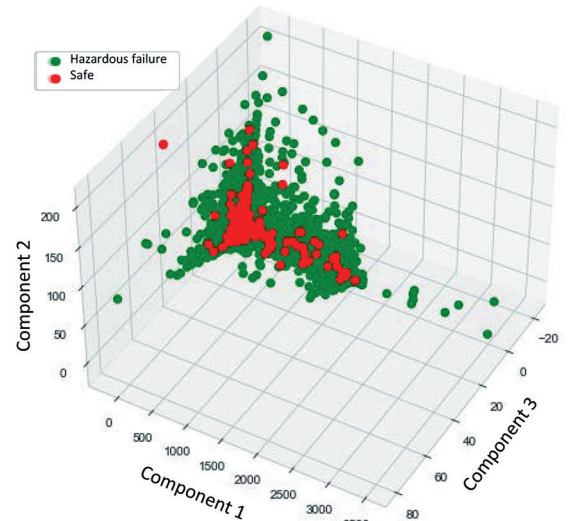


b) proportion of hazardous events month to month

Fig. 1. The proportion of hazardous events out of incidents involving RPSF



a) Class imbalance in a plane



b) Class imbalance in 3 dimensions

Fig. 2. Class imbalance at similar facilities

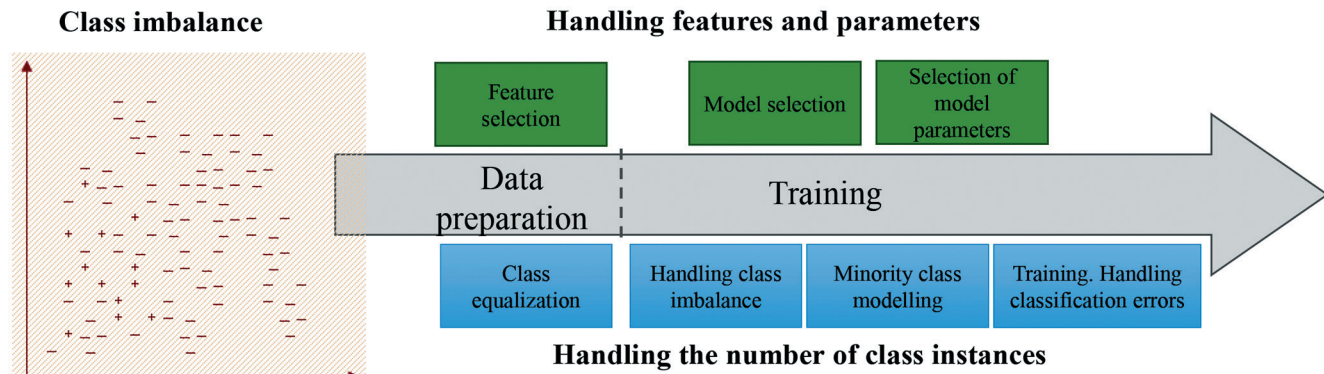


Fig. 3. Ways of improving the quality of classification

Data-level approach. SMOTE

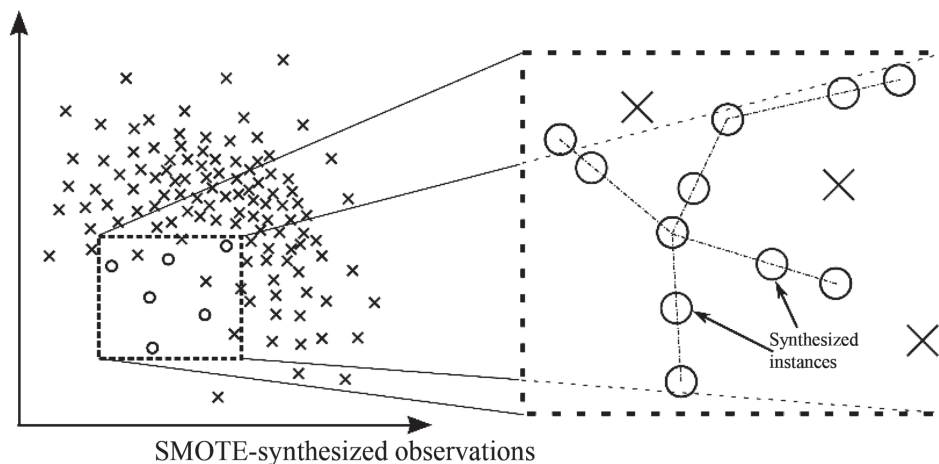
Fig. 4. Synthesized observation using *SMOTE*

Fig. 2a shows the ratio of classes on a plane with two coordinates, the number of incidents and the average number of days it takes to eliminate such incidents. It can be seen that there is no class boundary between the two features. Upon reducing the dimension using the dominant component analysis, let us move from a hundred of features to three (Fig. 2b). The graph better shows the concentration of hazardous events, yet there is still no clear boundary between classes.

Thus, improving the quality of incident classification is to involve handling class imbalance and selection of features. Otherwise, the quality of RPSF condition classification models will be unsatisfactory. (Table 1).

Table 1. Quality of classification models prior to the use of accuracy improvement techniques

Model	Accuracy of hazardous event prediction		Accuracy of safe state prediction	
	Training sample	Test sample	Training sample	Test sample
GBC	0	0.416	1	1
Logi	0	0	1	1
KNN	0	0	1	1
DTC	0	0.043	0.994	1

2. Ways of improving the accuracy of rare event prediction

The methods for improving the quality of classification can be divided into two groups: those based on features and parameters and those based on the number of class instances. Those methods can be used both at the stage of data preparation, and at the stage of training (Fig. 3).

The feature-based group includes:

- 1) methods of significant feature selection;
- 2) model selection and setting.

The methods based on the number of class instances include:

- 1) class equalization;
- 2) handling class imbalance;
- 3) minority class modelling;
- 4) selection of effective penalty function (handling classification error).

Although those methods are now widely known, there is no single algorithm of their combined application for the purpose of improved classification. The specificity and scope of data available in each particular situation force the researchers to experiment in search of an optimal combination of various methods. For instance, [1] proposed a technique that involves using a combination of classification

Table 2. Overview of the methods of handling class imbalance

Solution	Essential description of a group of methods	Advantage	Disadvantage
Data level approach			
MLSMOTE [7]	Pre-training stage: balancing by means of either undersampling, or oversampling in order to reduce the imbalance factor in the training data	Direct approach and wide application	Risk of overtraining
Diversified Sensitivity-based Undersampling (DSUS) [8]			
Ant Colony Optimization (ACO) [9]			
SMOTE [10]			
Evolutionary undersampling [11]			
Value (importance) boosting			
Cost-sensitive linguistic fuzzy rule [12]	Cost items (indicate the importance of class identification) designate the uneven importance of identification among classes. Increasing strategy may intentionally shift the training towards the classes associated with greater identification importance and ultimately improve the identification performance	A straightforward method, especially if the cost of error is known	Additional training costs due to finding an efficient cost matrix, especially when the real cost of error is unknown
Increasing cost sensitivity			
Feature selection			
Selection of minority class characteristics	Methods for selecting the features for the training	Helps solve class overlapping	Additional computational costs due to the requirement of data pre-processing
Density-based entity selection			
roc-based FAST generation of observations			
Correlation-based CFS generation of observations			
Algorithm-level approach			
Argument-based rule learning [14]	Specialized algorithms that specifically examine the distribution of class imbalance within datasets	Efficiency through modified algorithms for training solely based on the distribution of imbalance classes	It may be required to do pre-processing in order to balance-out uneven class distribution
Difference-based training [15]			
Fuzzy classifier [16]			
z-SVM [17]			
Hierarchical fuzzy rule [18]			
Distribution of conditional nearest class neighbour [19]			
k-NN sample generalization [20]			
Weighted nearest neighbour classifier [21]			
One-class training			
One-class training [22]	Classifier modelling on the minority class representation	Ease of use	Inefficient when used along with classification algorithms that are to be trained on the prevailing class
Class conditional nearest neighbour distribution (CCNND) [19]			
Economic training			
Bayesian SVM classifier [23]	A group of classification methods based on the cost of misclassification of both a false positive and a miss	A simple and quick processing technique	Inefficient if the actual cost is not available. Additional costs are introduced if cost investigation is required when the cost of error is unknown
Cost-sensitive SVM training [24]			
Cost-sensitive NN with PSO [25]			
SVM for adaptively asymmetric misclassification of cost [26]			
Ensemble-based method			
SMOTE and feature selection ensemble [27]	A group of methods based on the use of multiple classifiers that are trained directly on data and integration of the estimates for the purpose of developing a final classification solution	Multifaceted approaches	Complexity rises as the number of classifiers increases. Diversity is hard to achieve
GA ensembles [28]			
Ensembles for financial problems [29]			
Boosting in SVM ensemble [30]			
RUSBoost [31]			
SMOTEBoost [32]			

Solution	Essential description of a group of methods	Advantage	Disadvantage
Hybrid approach			
FTM-SVM	More than one machine learning algorithm is used in order to improve classification quality, often by combining them with other training algorithms for better results. Hybridization is used in order to simplify the sampling, selection of feature subset, optimization of the cost matrix and fine-tuning of the classical training algorithms	Symbiosis training through combination with other training algorithms is gaining popularity in class imbalance classification	A thorough project evaluation is required in order to take into account the differences between the used methods
F measure-based training [33]			
Linguistic fuzzy rule [34]			
Fuzzy classifier electronic algorithm [35]			
GA-based fuzzy rule extraction [36]			
Neuro-fuzzy [37]			
Neural network medical data [38]			
SMOTE neural networks [39]			
kNN classifier for medical data [40]			
NN trained with BP and PSO for medical data [41]			
Dependency tree kernels [42]			
Using cost sensitivity in trees [43]			
Undersampling and GA for SVM [44]			
Other methods			
ADASYN [45, 46]	Once the sample is created, it adds random small values to the points. In other words, instead of a linear correlation between the entire sample and the parent, they are a little more different	Adaptivity: the number of synthetic observations is based on the ratio of the majority to minority observations. ADASYN is focused on more complex data areas	The disadvantage of ADASYN is that it is easily affected by outliers
NearMiss [45, 47]	Random exclusion of majority class examples. When instances of two different classes are very close to each other, we remove the majority class instances in order to increase the space between the two classes. This helps the classification process	Can reduce sample overlapping between different classes	In case of undersampling, the number of insufficient samples cannot be controlled, while the mass samples that can be excluded are limited. Best used as a data cleansing method in combination with other methods
Edited nearest neighbours [45, 48]	Majority class instances are excluded if more than a half of their K neighbours do not belong to the majority class	Reduces misclassification of majority class instances	Cannot control the quantity of undersampling
Tomek Links [3, 45, 49]	Majority class instances are excluded that are in immediate proximity to minority class instances	All items that are immediate neighbours belong to the same class that can be better classified	
Neighbourhood cleaning rule [45]	This strategy also aims to remove those examples that negatively affect the outcome of minority class classification. For that purpose, all examples are classified according to the rule of the three nearest neighbours	Improves the accuracy of minority class instance classification	May result in a sample size that would not be sufficient for training
Condensed nearest neighbour [45, 50]	The method allows finding differences between similar examples that belong to different classes	Finding differences between similar examples that belong to different classes	

algorithms and the *RFE*, *Random Forest* and *Boruta* methods of feature selection with preliminary class balancing by means of *SMOTE* and *ADASYN* random sampling. The paper shows an increase in classification accuracy up to 98% (from 93%). However, study [2] dedicated to text classification shows that rather than trying to modify the distribution, it would be more efficient to work with decision threshold modification and the weights of errors of various kinds. Additionally, the author introduces a separation of unbalanced data into moderately unbalanced data (class ratio 7 to 1) and strongly unbalanced data (class ratio 14 to 1). Paper [2] showed that handling class ratios using strongly unbalanced data in case of some models causes improved classification accuracy. While classifying user requests, [3] identified that using class-balancing methods may not only fail to provide any results, but also cause reduced classification accuracy. Contrary to [3], [4] demonstrated the efficiency of data sampling methods.

The purpose of this paper is to practically demonstrate the effect of methods of handling class instances on the classification accuracy of hazardous events affecting RPSF.

2.1. Overview of the methods of handling class imbalance

One of the primary methods of handling class imbalance is *SMOTE*, whose algorithm was developed in early 2002 [5]. Currently, there are a number of modifications of this method. It also inspired the development of other algorithms of handling unbalanced data. *SMOTE* is at the top of the group of data sampling methods, i.e., those that involve increasing the number of minority class instances. Its basic principle is shown in Fig. 4.

Synthetic instances are generated in the “function space” rather than the “data space”. Minority class samples are replenished by introducing synthetic examples along segments of the line that connect any/all of the nearest neighbours of the minority class k [5]. *SMOTE* is an oversampling method.

Another way of handling unbalanced data is to reduce the number of the majority class instances (undersampling or reindexing).

Combinations of various sampling strategies constitute hybrid methods that involve sequential application of over-sampling and undersampling algorithms. A visualization of the processes and results of the above strategies is shown in Fig. 5.

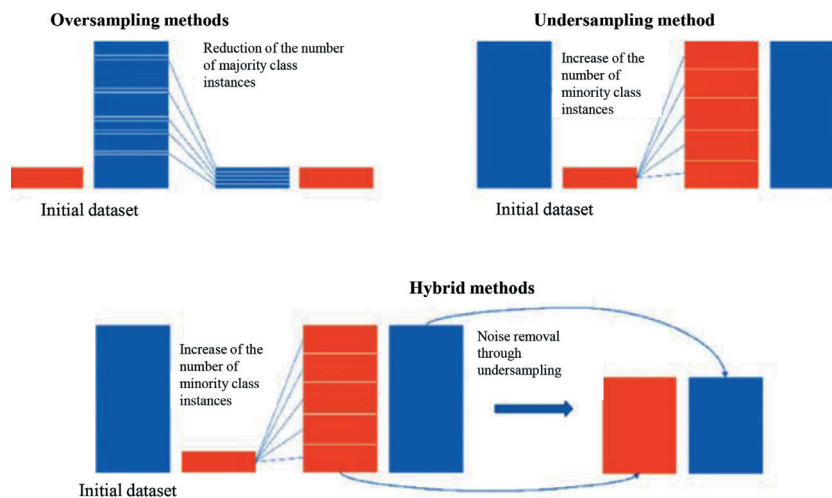
[6] made an overview of the methods of handling class imbalance. Table 2 shows information on them along with the currently-employed methods.

Based on the advantages and disadvantages of the methods examined in Table 2, as well as the experience of the researchers studying data sampling techniques [51, 52, 53, 54], *NearMiss* [47, 55] was chosen as the primary method of improving the classification of unbalanced data. The purpose of *NearMiss* is to balance the distribution of observations across minority and majority classes by estimating the distance between instances from different classes. The *NearMiss* implementation of the *imblearn Python* library (includes a set of tools for unbalanced datasets in machine learning) involves three strategies of class instance selection:

Strategy 1 (*version* = 1). Selection of observations out of the majority class, for which the average distance to k nearest observations out of the minority class is the smallest (by default $k = 3$, training variable). Only those observations without hazardous events will be retained that are the nearest to those with hazardous events;

Strategy 2 (*version* = 2). Selection of observations from the majority class, for which the average distance to k furthest observations of the minority class is the smallest (by default, $k = 3$, training variable). Only those observations without hazardous events will be retained that are at the centre of the mass of the intersection of sets of majority and minority classes;

Strategy 3 (*version* = 3). First, for each observation out of the majority class, M nearest neighbours will be retained (by default, $M = 3$, training variable). Then, observations out of the minority class are selected, for which the average



The operating principle of hybrid methods in respect to the problem of class balancing

Fig. 5. Data sampling strategies

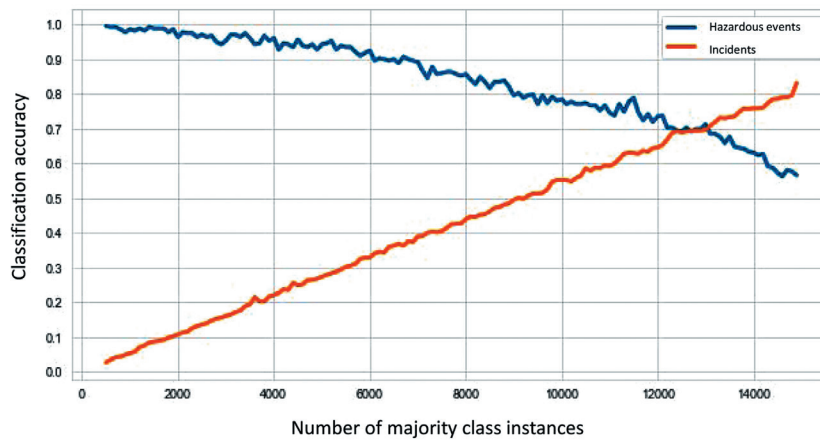


Fig. 6. Classification accuracy on main sample, GBC, NearMiss 2

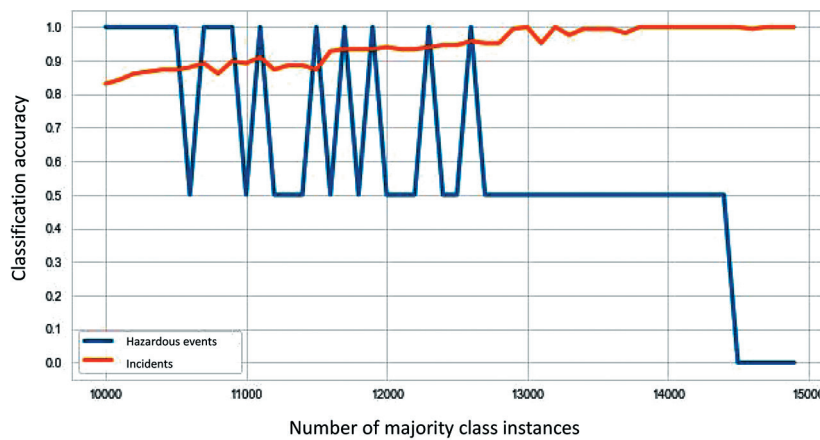


Fig. 7. Accuracy on validation sample, GBC, NearMiss 1

distance to N nearest neighbours is the longest (by default, $N = 3$, training variable).

Unlike in [1-4, 52, 53], the focus of our attention is not the accuracy or classification error indicator, but rather the detail of the hazardous event prediction accuracy and the incident prediction accuracy. Additionally, due to the fact that 98% of observations are in the majority class, it should be expected that the accuracy of most trained classifiers will be 98%. In this context, it is easy to conclude that class balancing is inefficient, as some researchers do.

In the course of the study, a combination of several methods of improving the quality of unbalanced data classification was used, namely *GridSearchCV* and *NearMiss*, as well as various quality functions.

The parameters that vary in the course of the experiment:

1) within the *GridSearchCV* function:

- (1) GBS: *random_state*, *tol*, *max_depth*;
- (2) Logi: *tol*, *class_weight*, *max_iter*, *solver*, *random_state*, *C*;
- (3) KNN: *n_neighbors*, *weights*, *metric*;
- (4) DTC: *criterion*, *max_depth*, *min_samples_split*, *max_features*;

2) quality functions: *max_error*, *balanced_accuracy*, *accuracy*, *neg_log_loss*, *explained_variance*, *neg_mean_squared_error*, *neg_mean_squared_log_error*, *neg_median_absolute_error*, *r2*;

3) NearMiss parameters: instance selection strategies, number of minority class instances, number of majority class instances.

Diagram of the experiment:

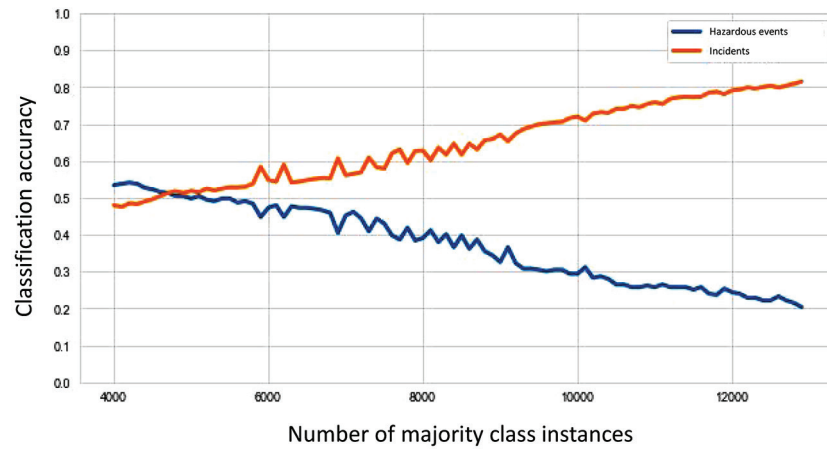
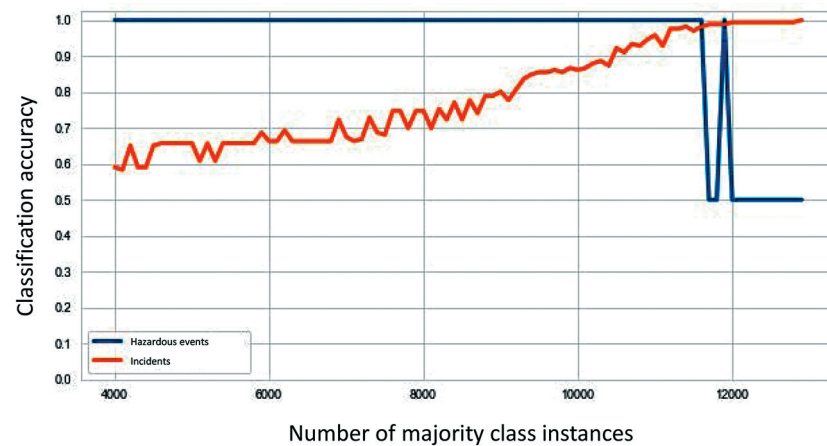
1) division of the sample into the master and the validation. The master sample includes observations collected before 2019, while the validation includes those collected after 2019;

2) division of the master sample into the training and test. The method is *train_test_split*, the test sample size is 20%.

The experiments have established that the best accuracy is achieved under the *neg_log_loss* quality function. Due to the extremely small number of minority class instances (less than 2%) the quality of this indicator's variation has no impact on the classification quality. In the course of most experiments the number of minority class instances was fixed and equalled the maximum possible.

Let us examine the experimental results. Fig. 6 shows the graph of accuracy of a GBC event prediction depending on the number of majority class instances on the NearMiss 2 master sample.

The initial accuracy characteristics (without NearMiss) are: accuracy on the test part of the training sample: 0.983; accuracy of hazardous event prediction: 0.416; accuracy of non-hazardous incident prediction: 1. As can be seen in Fig. 6, there is a point where the accuracy graphs of hazard-

Fig. 8. Accuracy on main sample, *Logi*, *NearMiss 1*Fig. 9. Accuracy on validation sample, *Logi*, *NearMiss 1*.

ous events and incidents classification intersect. Essentially, the higher this point is, the more accurate the final model will be. Similar graphs were obtained for the *NearMiss 1* sampling strategy. An accuracy of 70% appears to be a good result in case of the starting accuracy of hazardous event classification of 41.6%. The following experiment was conducted under selection strategy no. 2. For clarity, Fig. 7 only shows the intersection of the graphs that characterize the validation sample.

The accuracy of the classifier on the validation sample (v-sample) proved to be significantly higher than that on the master sample. According to the analysis, that was due to the size of the v-sample (the classifier “did have enough time” to make many mistakes), as well as the months covered by the prediction. The accuracy over periods similar to those of the v-sample in the master sample proved to be higher than average.

Before proceeding to the analysis of other data, let us revisit the topic of using the classifier’s accuracy indicator as an efficiency characteristic. The GBC classifier accuracy on unbalanced data was 0.983. The accuracy of hazardous event prediction on the v-sample is 0. On balanced data the accuracy was 0.9811 (lower than on the original sample size), while the hazardous event prediction accuracy on the v-sample was 1.0, the accuracy of incident prediction was

0.934. Similar figures were obtained for other classifiers. Thus, when studying the classification accuracy of unbalanced data, one cannot rely on one “convoluted” quality indicator.

Fig. 8 and 9 show the accuracy graphs of the *Logi* model for the master sample and the *NearMiss 1* v-sample.

The graphs in Fig. 8 and 9 show that the best accuracy indicators for the master and v-sample are achieved under

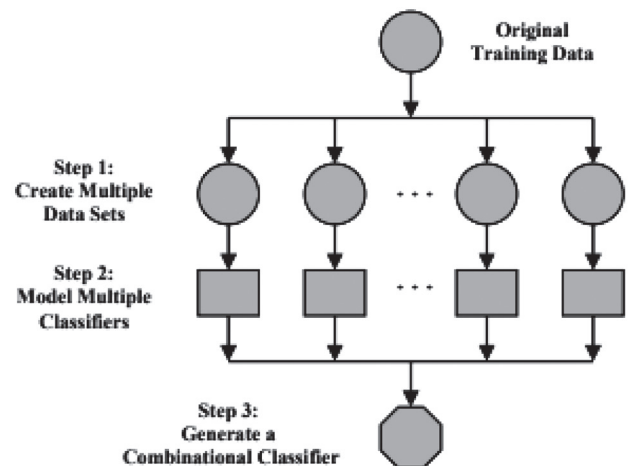


Fig. 10. General algorithm of ensemble methods

different class balances. That suggests that the models rejected on the basis of the master sample data may produce better results on new data. One of the ways of solution that problem is to use ensemble methods [56]. The key idea of studying ensembles of classifiers is to build several classifiers from the original data and then to aggregate the predictions when classifying unknown samples (Fig. 10). The application of such methods may be the subject matter of further research.

Analysis and conclusions

1. As safety systems evolve, the number of hazardous events decreases, yet the cost of potential consequences grows. Out of RPSF incidents, the proportion of hazardous events does not exceed 2% per year. For the purpose of improving the accuracy of classification and prediction of event types in case of class imbalance, sample balancing methods are to be used.

2. More than 50 methods are currently in active use that allow handling unbalanced samples. However, the publications known to the authors do not address the matter of analysing simultaneous changes in the prediction accuracy of minority and majority class instances. The paper presents graphs of classification accuracy of RPSF incidents on the training and validation sample.

3. A method of handling unbalanced data is proposed that includes a combination of several ways of improving the quality of unbalanced data classification: model parameter setting, selection of the indicator of quality and ratio of the minority to the majority class instance number using *NearMiss*.

4. The methods of dealing with class imbalance allows significantly increasing the accuracy of predicting minority class instances (hazardous events) from 0 to 70-90%. However, as the accuracy of prediction of rare events increases, the accuracy of prediction of minority class instances decreases.

5. The study may pave the way for the application of hybrid methods of classifying and predicting events, as well as for the development of a metric of training quality based on the characteristic of the intersection point of classification accuracy graphs of the instances of different classes.

References

1. Sevastianov L.A., Shetinin E.Yu. On methods for improving the accuracy of multiclass classification on imbalanced data. *Informatics and applications* 2020;14(1):63-70. (in Russ.)
2. Sadov M.A. Study of the methods of text classification for unbalanced data. *Polymathis Scientific Journal* 2016;2:28-41. (in Russ.)
3. Maslikhov S.R., Mokhov A.S., Tolcheev V.Yu. [Building balanced classes in respect to user query classification]. [Proceedings of the 5th International Science and Practice Conference Remote Education Technologies]; 2020. P. 245-248. (in Russ.)
4. Shipitsyn A.V., Zhuravleva N.V. Evaluation of online mortgage applications with machine learning algorithms. *Herald of the Belgorod University of Cooperation, Economics and Law* 2016;4(60):199-209. (in Russ.)
5. Chawla N.V., Bowyer W.B., Hall L.O. et al. Smart: Synthetic Minority Over-sampling Technique. *Journal of Artificial Intelligence Research* 2002;16:321-357.
6. Ali A., Shamsuddin S.M., Ralescu A. Classification with class impact problem: a review. *International Journal of Advances in Soft Computing* 2013;7:176-204.
7. Mladenic D., Grobelnik M. Feature selection for unbalanced class distribution and national scores. Proceedings of the Sixteenth International Conference on Machine Learning (ICML 1999). Bled (Slovenia); 1999. p. 258-267.
8. Yang T.-N., Wang S.-D. Robust algorithms for principal component analysis. *Pattern Recognition Letters* 1999;20(9):927-933.
9. Yu H., Ni J., Zhao J. ACOSampling: An ant colony optimization-based undersampling method for classifying imbalanced DNA microarray data. *Neurocomputing* 2013;101(0):309-318.
10. Chawla N.V. SMOTE: synthetic minority over-sampling technique. arXiv:1106.1813; 2002.
11. Garcia S., Herrera F. Evolutionary undersampling for classification with imbalanced datasets: Proposals and taxonomy. *Evolutionary Computation* 2009;17(3):275-306.
12. Yin L. Feature selection for high-dimensional imbalanced data. *Neurocomputing* 2013;105(0):3-11.
13. Sun Y. Cost-sensitive boosting for classification of imbalanced data. *Pattern Recognition* 2007;40(12):3358-3378.
14. Luukka P. Nonlinear fuzzy robust PCA algorithms and similarity classifier in bankruptcy analysis. *Expert Systems with Applications* 2010;37(12):8296-8302.
15. Zheng Z., Wu X., Srihari R. Feature selection for text categorization on imbalanced data. *ACM SIGKDD Explorations Newsletter* 2004;6(1):80-89.
16. Visa S., Ralescu A.L. Fuzzy Classifiers for Imbalanced Data Sets. University of Cincinnati, Computer Science Dept. Cincinnati (OH, United States); 2007.
17. Imam T., Ting K., Kamruzzaman J. z-SVM: An SVM for Improved Classification of Imbalanced Data. AI 2006: Advances in Artificial Intelligence. 19th Australian Joint Conference on Artificial Intelligence. Hobart (Australia); 2006. p. 264-273.
18. Fernández A., M.J. del Jesus, Herrera F. Hierarchical fuzzy rule based classification systems with genetic rule selection for imbalanced data-sets. *International Journal of Approximate Reasoning* 2009;50(3):561-577.
19. Kriminger E., Principe J.C., Lakshminarayan C. Nearest Neighbor Distributions for imbalanced classification. The 2012 international joint conference on neural networks (IJCNN). Brisbane (QLD, Australia); 2012. p. 1-5.
20. Li Y., Zhang X. Improving k nearest neighbor with exemplar generalization for imbalanced classification. In: Proceedings of Advances in knowledge discovery and data mining: 15th Pacific-Asia Conference, Part II. Shenzhen (China); 2011. p. 321-332.

21. Candès E.J. Robust principal component analysis. *Journal of the ACM (JACM)* 2011;58(3):11.
22. Japkowicz N., Myers C., Gluck M. A novelty detection approach to classification. *IJCAI* 1995;1:518–523.
23. Jolliffe I. Principal component analysis. Encyclopedia of Statistics in Behavioral Science. John Wiley & Sons, Ltd; 2005.
24. Cao P., Zhao D., Zaiane O. An Optimized Cost-Sensitive SVM for Imbalanced Data Learning. In: Proceedings of Advances in Knowledge Discovery and Data Mining: 17th Pacific-Asia Conference. Part II. Gold Coast (Australia); 2013. p. 280-292.
25. Cao P., Zhao D., Zaiane O. A PSO-based Cost-Sensitive Neural Network for Imbalanced Data Classification. In: Revised Selected Papers of Trends and Applications in Knowledge Discovery and Data Mining. International Workshops: DMAPs, DANTH, QIMIE, BDM, CDA, CloudSD. Gold Coast (Australia); 2013. p. 452-463.
26. Wang X., Shao H., Japkowicz N., et al. Using SVM with Adaptive Asymmetric Miseducation Costs for Mine-like objects Detection. In: Proceedings of the 11th International Conference on Machine Learning and Applications. Boca Raton (Florida, USA); 2012. p. 78-82.
27. Yang P., Liu W., Zhou B.B. et al. Ensemble-based wrapper methods for feature selection and class imbalance learning. In: Proceedings of Advances in Knowledge Discovery and Data Mining: 17th Pacific-Asia Conference, Part I. Gold Coast (Australia); 2013. p. 544-555.
28. Yu E., Cho S. Ensemble based on GA wrapper feature selection. *Computers & Industrial Engineering* 2006;51(1):111-116.
29. Liao J.-J. An ensemble-based model for two-class imbalanced financial problem. *Economic Modelling* 2014;37(0):175-183.
30. Liu Y., An A., Huang X. Boosting prediction accuracy on imbalanced datasets with SVM ensembles. In: Proceedings of Advances in Knowledge Discovery and Data Mining, 10th Pacific-Asia Conference. Singapore; 2006. p. 107-118.
31. Seiffert C. RUSBoost: A hybrid approach to alleviating class imbalance. *Systems, Man and Cybernetics. Part A: Systems and Humans. IEEE Transactions* 2010;40(1):185-197.
32. Chawla N.V. SMOTEBoost: Improving prediction of the minority class in boosting. In: Proceedings of Knowledge Discovery in Databases: PKDD 2003, 7th European Conference on Principles and Practice of Knowledge Discovery in Databases. Cavtat-Dubrovnik (Croatia); 2003. p. 107-119.
33. Wasikowski M., Chen X.-W. Integrating the small sample class impact problem using feature selection. *Knowledge and Data Engineering. IEEE transactions* 2010;22(10):P.1388-1400.
34. Martino M.D. Novel Classifier Scheme for Unbalance Problems. *Pattern Recognition Letters* 2013;34(10):1146–1151.
35. Fernández A. A study of the behaviour of linguistic fuzzy rule based classification systems in the framework of imbalanced data-sets. *Fuzzy Sets and Systems* 2008;159(18):2378-2398.
36. Le X., Mo-Yuen C., Taylor L.S. Power Distribution Fault Cause Identification With Imbalanced Data Using the Data Mining-Based Fuzzy Classification E-Algorithm. *Power Systems. IEEE Transactions* 2007;22(1):164-171.
37. Soler V. Imbalanced Datasets Classification by Fuzzy Rule Extraction and Genetic Algorithms. Data Mining Workshops. ICDM Workshops. Sixth IEEE International Conference; 2006.
38. Hung C.-M., Huang Y.-M. Conflict-sensitivity contexture learning algorithm for mining interesting patterns using neuro-fuzzy network with decision rules. *Expert Systems with Applications* 2008;34(1);159-172.
39. Jeatrakul P., Wong K.W., Fung C.C. Classification of imbalanced data by combining the complementary neural network and SMOTE algorithm. Proceedings of the 17th International Conference on Neural Information Processing: Models and Applications. Part II. Sydney (Australia); Springer-Verlag. p. 152-159.
40. Malof J.M., Mazurowski M.A., Tourassi G.D. The effect of class imbalance on case selection for case-based classifiers: An empirical study in the context of medical decision support. *Neural Networks* 2012;25(0):141-145.
41. Mazurowski M.A. Training neural network classifiers for medical decision making: the effects of imbalanced datasets on classification performance. *Neural networks* 2008;21(2-3):427-436.
42. Culotta A., Sorensen J. Dependency tree kernels for relation extraction. In: Proceedings of the 42nd Annual Meeting on Association for Computational Linguistics. Association for Computational Linguistics; 2004.
43. Drummond C., Holte R.C. Exploiting the cost (in) sensitivity of decision tree splitting criteria. *ICML*; 2000.
44. Al-Shahib A., Breitling R., Gilbert D. Feature selection and the class imbalance problem in predicting protein function from sequence. *Applied Bioinformatics* 2005;4(3):195-203.
45. Koziarski M. Radial-Based Undersampling for imbalanced data classification. *Pattern Recognition* 2020;102.
46. He H., Bai Y., Garcia E.A. et al. ADASYN: adaptive synthetic sampling approach for imbalanced learning. In: Proceedings of IEEE International Joint Conference on Neural Networks (IEEE World Congress on Computational Intelligence); 2008. p. 1322-1328.
47. Mani I., Zhang I. kNN approach to unbalanced data distributions: a case study involving information extraction. In: Proceedings of Workshop on Learning from Imbalanced Datasets. Vol. 126; 2003.
48. Wilson D.L. Asymptotic properties of nearest neighbor rules using edited data. *IEEE Trans. Syst. Man Cybern.* 1972;2(3): 408-421.
49. Tomek I. Two modifications of CNN. *IEEE Trans. Syst. Man Cybern.* 1976;6:769-772.
50. Hart P. The condensed nearest neighbor rule. *IEEE Trans. Inf. Theory* 1968;14(3):515-516.

51. Makhsotova Ts.V. [Study of classification methods in the case of class imbalance]. *Science Magazine* 2017;5(18). (accessed July 5, 2020). Available at: <https://cyberleninka.ru/article/n/issledovanie-metodov-klassifikatsii-pri-nesbalansirovannosti-klassov>. (in Russ.)

52. Kavrin D.A., Subbotin S.A. The methods for quantitative solving the class imbalance problem. *Radio Electronics, Computer Science, Control* 2018;1. (accessed July 6, 2020). Available at: <https://cyberleninka.ru/article/n/metody-kolichestvennogo-resheniya-problemy-nesbalansirovannosti-klassov>. (in Russ.)

53. Yi L., Hong G., Feldkamp L. Robust neural learning from unbalanced data samples. In: Proceedings of IEEE International Joint Conference on Neural Networks. IEEE World Congress on Computational Intelligence (Cat. No.98CH36227). Vol. 3. Anchorage (USA); 1998. p. 1816-1821.

54. Al-Stouhi S., Reddy C.K. Transfer learning for class imbalance problems with inadequate data. *Knowledge and Information Systems* 2016;48:201-228.

55. Near-Miss – version 0.9.0.dev0. API reference. (accessed July 10, 2021). Available at: https://imbalanced-learn.org/dev/references/generated/imblearn.under_sampling.NearMiss.html.

56. Sun Y. Cost-Sensitive Boosting for Classification of Imbalanced Data. *Pattern Recognition* 2007;40(12):3358-3378.

About the authors

Olga B. Pronevich, Head of Unit, JSC NIIAS, 27, bldg 1 Nizhegorodskaya St., Moscow, 109029, Russian Federation, phone: +7 (495) 786 68 57; e-mail: obpronevich@gmail.com.

Mikhail V. Zaytsev, Lead Specialist, JSC NIIAS, 27, bldg 1 Nizhegorodskaya St., Moscow, 109029, Russian Federation, phone: +7 (495) 786 68 57; e-mail: m.v.zaicev@mail.ru.

The authors' contribution

Pronevich O.B. investigated the problem of class imbalance when predicting the occurrence of hazardous events at railway power supply facilities, analysed the existing methods of unbalanced data processing, as well as conducted a series of experiments for evaluating the effect of different proportions of the instances of the minority and majority classes. The author also defined a vision for further research.

Zaytsev M.V. analysed widely used methods of event classification under class imbalance, studied the effect of various strategies of *NearMiss* selection on the quality of hazardous event classification in railway transportation.

Conflict of interests

The authors declare the absence of a conflict of interests.

Dear colleagues,

The COVID-19 pandemic, the greatest disease of our time that has afflicted the entire world, has also taken its toll on the Editorial Board of the Dependability Journal. On July 18, 2021, it took the life of our colleague, friend, the Journals' Executive Editor Dr. Alexey Zamyshliaev. He was only 41. Alexey's family, friends and colleagues were stunned by his sudden demise. It is hard to conceive that this life-loving, intelligent, kind man is no longer with us. We lost a wonderful man, a friend always ready to help anyone around. He was always respectful and receptive to the advice of his mentors. He was the prop and stay of his family, a support for the people around. He took great care of his mother's health and well-being. He encompassed his spouse and children with great care and attention. He loved life and was loved by everyone who lived and worked by his side.

Alexey was a talented scientist. He became Doctor of Engineering at the age of 34. He authored a great number of publications, research papers and books. He devoted a lot of time to the development of the Dependability Journal, its promotion abroad. He was a member of the Asset Management Working Group of the International Union of Railways. A great organiser, he managed to unite hundreds of researchers and experts into a single, tightly-welded team capable of solving the problems associated with the deployment of advanced digital technology in railway transportation, including computer simulation, intelligent control and management systems, AI-based physical railway asset management systems, etc.

Alexey will be always remembered by the managers and employees of the Russian Railways as a brilliant expert who solved complex problems of railway transportation management. He led the design and develop-



ment of a number of systems that now help manage traffic safety, technical assets and in-station processes. Under his leadership, the railways' essential information systems were created: JSC RZD's Automated System for Station Technological Instructions Management (AS TRA), Automated Traffic Safety Management System (AS RB), Integrated Automated System for Tracking, Supervision and Elimination of Technical Failures and Dependability Analysis (KAS-ANT), Integrated Automated System for Recording, Investigation and Analysis of Process Violations (KASAT), Single Corporate Platform for Management of Resources, Risks and Dependability (URRAN), Automated Management System of JSC RZD Situation Centre and many others.

He passed away full of ideas and plans. He wanted to bring the knowledge and experience he had accumulated to the railways of other countries, as well as to related industries. He planned to teach students. He loved to live an interesting, colourful and active life. He loved his family and his life.

He will always be missed by all of us who were his friends, colleagues and acquaintances.

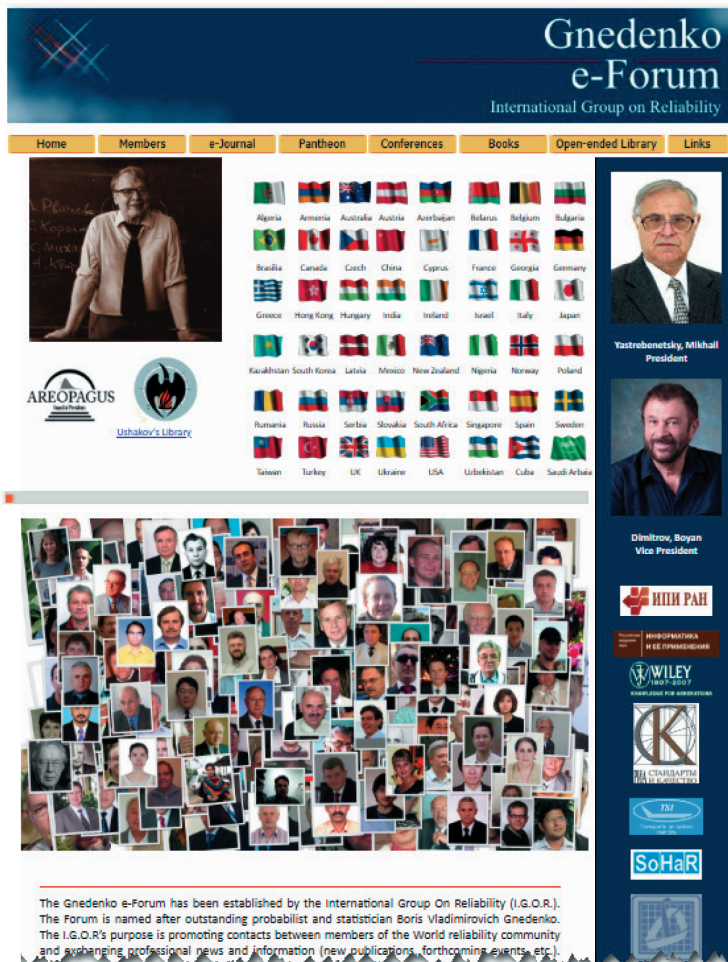
The memory of Alexey Zamyshliaev will always dwell with us.

Editorial Board, Dependability Journal



GNEDENKO FORUM

INTERNATIONAL GROUP ON RELIABILITY



The Gnedenko Forum was founded in 2004 by an unofficial international group of experts in the dependability theory for the purpose of professional support of researches from all over the world who are interested in studying and developing the scientific, technical and other aspects of the dependability theory, risk analysis and safety in the theoretical and practical domains.

The Forum exists on the Internet as a non-for-profit organization. It aims to involve into joint discussion and communication technical experts interested in developing the dependability theory, safety and risk analysis regardless of their home country and membership in whichever organization.

The Forum acts as an impartial and neutral entity that delivers scientific information to the press and public as regards the matters of safety, risk analysis and dependability of complex technical systems. It publishes reviews, technical documents, technical reports and research essays for the purpose of dissemination of knowledge and information.

The Forum is named after Boris V. Gnedenko, an outstanding Soviet mathematician, expert in the probability theory and its applications, member of the Ukrainian Academy of Sciences. The Forum is the platform for distribution of information on educational grants, academic and professional positions related to dependability, safety and risk analysis all over the world.

Currently, the Forum has 500 members from 47 countries.

Since January 2006, the Forum has been publishing its quarterly journal, Reliability: Theory & Applications (www.gnedenko.net/RTA). The Journal is registered in the Library of Congress (ISSN 1932-2321) and publishes articles, reviews, memories, information and literature references regarding the theory and application of dependability, survivability, maintenance, risk analysis and management methods.

Since 2000, the Journal is indexed in Scopus.



Membership in the Gnedenko Forum does not imply any obligations. It is only required to send your photograph and a brief professional biography (resume) to a.bochkov@gmail.com. Templates can be found at <http://www.gnedenko.net/personalities.htm>.

www.gnedenko.net

DEPENDABILITY JOURNAL ARTICLE SUBMISSION GUIDELINES

Article formatting requirements

Articles must be submitted to the editorial office in electronic form as a Microsoft Office Word file (*.doc or *.docx extension). The text must be in black, on a A4 sheet with the following margins: 2 cm for the left, top and bottom margins; 1.5 or 2 cm for the right margin. An article cannot be shorter than 5 pages and longer than 12 pages (can be extended upon agreement with the editorial office). The article is to include the structural elements described below.

Structure of the article

The following structural elements must be separated with an *empty line*. Examples of how they must look in the text are shown *in blue*.

1) Title of the article

The title of the article is given in the English language.

Presentation: The title must be in 12-point Times New Roman, with 1.5 line spacing, fully justified, with no indentation on the left. The font face must be bold. The title is not followed by a full stop.

An example:

Improving the dependability of electronic components

2) Author(s)' name.

This structural element for each author includes:

In English: second name and first name as "First name, Second name" (John Johnson).

Presentation: The authors' names must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be bold. The authors' names are separated with a comma. The line is not followed by a full stop.

An example:

John Johnson¹, Karen Smith^{2*}

3) The author(s)' place of employment

The authors' place of employment is given in English. Before the place of employment, the superscripted number of the respective reference to the author's name is written.

Presentation: The reference to the place of employment must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal. Each place of employment is written in a new line. The lines are not followed by a full stop.

An example:

¹ Moscow State University, Russian Federation, Moscow

² Saint Petersburg Institute of Heat Power Engineering, Russian Federation, Saint Petersburg

4) The e-mail address of the author responsible for maintaining correspondence with the editorial office

Presentation: The address must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal, all symbols must be lower-case. Before the address reference, symbol * is written. The title is not followed by a full stop.

An example:

*johnson_j@aaa.net

5) Abstract of the article

This structural element includes a structured summary of the article with the minimal size of 350 words and maximum size of 400 words. The abstract is given in the English language. The abstract must include (preferably explicitly) the following sections: Aim; Methods; Results/Findings; Conclusions. The abstract of the article should not include newly introduced terms, abbreviations (unless universally accepted), references to literature.

Presentation: The abstract must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal, except "Abstract", "Aim", "Methods", "Conclusions", that (along with the full stop) must be in bold. The text of the abstract must not be paragraphed (written in a single paragraph).

An example:

Abstract. Aim. Proposing an approach ... taking into consideration the current methods. **Methods.** The paper uses methods of mathematical analysis, ..., probability theory. **Results.** The following findings were obtained using the proposed method ... **Conclusion.** The approach proposed in the paper allows...

6) Keywords

5 to 7 words associated with the paper's subject matter must be listed. It is advisable that the keywords complemented the abstract and title of the article. The keywords are written in English.

Presentation: The text must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal, except “**Keywords:**” that (along with the colon) must be in bold. The text must not be paragraphed (written in a single paragraph). The text must be followed by a full stop.

An example:

Keywords: dependability, functional safety, technical systems, risk management, operational efficiency.

7) Text of the article

It is recommended to structure the text of the article in the following sections: Introduction, Overview of the sources, Methods, Results, Discussion, Conclusions. Figures and tables are included in the text of the article (the figures must be “In line with text”, not “behind text” or “in front of text”; not “With Text Wrapping”).

Presentation:

The titles of the sections must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be bold. The titles of the sections (except the Introduction and Conclusions) may be numbered in Arabic figures with a full stop after the number of a section. The number with a full stop must be separated from the title with a no-break space (Ctrl+Shift+Spacebar).

The text of the sections must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with a 1.25-cm indent. The font face must be normal. The text of the sections must be paragraphed. There must be no indent in the paragraph that follows a formula and contain notes to such formula, e.g.:

where n is the number of products.

An example:

1. State of the art of improving the dependability of electronic components

An analysis of Russian and foreign literature on the topic of this study has shown that ...

Figures (photographs, screenshots) must be of good quality, suitable for printing. The resolution must be at least 300 dpi. If a figure is a diagram, drawing, etc. it should be inserted into the text in editable form (Microsoft Visio). All figures must be captioned. Figures are numbered in Arabic figures in the order of their appearance in the text. If a text has one figure, it is not numbered. References to figures must be written as follows: “Fig. 3. shows that ...” or “It is shown that ... (see. Fig. 3.)” The abbreviation “Fig.” and number of the figure (if any) are always separated with a no-break space (Ctrl+Shift+Spacebar). The caption must include the counting number of the figure and its title. It must be placed a line below the figure and center justified:

Fig. 2. Description of vital process

Captions are not followed by a full stop. *With center justification there must be no indent!* All designations shown in figures must be explained in the main text or the captions. The designations in the text and the figure must be identical (including the differences between the upright and oblique fonts). *In case of difficulties with in-text figure formatting, the authors must – at the editorial office’s request – provide such figures in a graphics format (files with the *.tiff, *.png, *.gif, *.jpg, *.eps extensions).*

The tables must be of good quality, suitable for printing. The tables must be editable (not scanned or in image format). All tables must be titled. Tables are numbered in Arabic figures in the order of their appearance in the text. If a text has one table, it is not numbered. References to tables must be written as follows: “Tab. 3. shows that ...” or “It is shown that ... (see. tab. 3.)” The abbreviation “tab.” and number of the table (if any) must be always separated with a no-break space (Ctrl+Shift+Spacebar). The title of a table must include the counting number and its title. It is placed a line above the table with center justification:

Table 2. Description of vital process

The title of a table is not followed by a full stop. *With center justification there must be no indent!* All designations featured in tables must be explained in the main text. The designations in the text and tables must be identical (including the differences between the upright and oblique fonts).

Mathematical notations in the text must be written in capital and lower-case letters of the Latin and Greek alphabets. Latin symbols must always be oblique, except function designators, such as sin, cos, max, min, etc., that must be written in an upright font. Greek symbols must always be written in an upright font. The font size of the main text and mathematical notations (including formulas) must be identical; in Microsoft Word upper and lower indices are scaled automatically.

Formulas may be added directly into the text, for instance:

Let $y = a \cdot x + b$, then ...,

or written in a separate line with center justification, e.g.:

$$y = a \cdot x + b.$$

In formulas both in the text, and in separate lines, the punctuation must be according to the normal rules, i.e. if a formula concludes a sentence, it is followed by a full stop; if the sentence continues after a formula, it is followed by a comma (or no punctuation mark). In order to separate formulas from the text, it is recommended to set the spacing for the formula line 6 points before and 6 points after). If a formula is referenced in the text of an article, such formula must be written in a separate line with the number of the formula written by the right edge in round brackets, for instance:

$$y = a \cdot x + b. \quad (1)$$

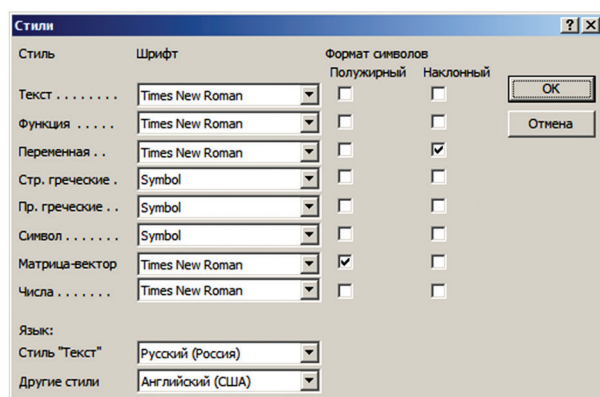
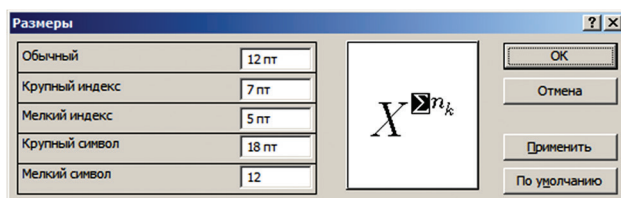
If a formula is written in a separate line and has a number, such line must be right justified, and the formula and its number must be tab-separated; tab position (in cm) is to be chosen in such a way as to place the formula roughly at the center. Formulas that are referenced in the text must be numbered in Arabic figures in the order of their appearance in the text.

Simple formulas should be written without using formula editors (in MS Word, Latin should be used, as well as the “Insert” menu + “Special Characters”, if Greek letters and mathematical operators are required), while observing the required slope for Latin symbols, for example:

$$\Omega = a + b \cdot \theta.$$

If a formula is written without using a formula editor, letters and +, -, = signs must be separated with no-break spaces (Ctrl+Shift+Spacebar).

Complex formulas must be written using a formula editor. In order to avoid problems when editing and formatting formulas it is highly recommended to use Microsoft Equation 3.0 or MathType 6.x. In order to ensure correct formula input (symbol size, slope, etc.), below are given the recommended editor settings.



When writing formulas in an editor, if brackets are required, those from the formula editor should be used and not typed on the keyboard (to ensure correct bracket height depending on the formula contents), for example (Equation 3.0):

$$Z = \frac{a \cdot \left(\sum_{i=1}^n x_i + \sum_{j=1}^m y_j \right)}{n + m} \quad (2)$$

Footnotes in the text are numbered with Arabic figures, placed page by page. Footnotes may include: references to anonymous sources on the Internet, textbooks, study guides, standards, information from websites, statistic reports, publications in newspapers, magazines, autoabstracts, dissertations (if the articles published as the result of thesis research cannot be quoted), the author's comments.

References to bibliographic sources are written in the text in square brackets, and the sources are listed in the order of citation (end references). The page number is given within the brackets, separated with a comma and a space, after the source number: [6, p. 8].

8) Acknowledgements

This section contains the mentions of all sources of funds for the study, as well as acknowledgements to people who took part in the article preparation, but are not among the authors. Participation in the article preparation implies: recommendations regarding improvements to the study, provision of premises for research, institutional supervision, financial support, individual analytical operations, provision of reagents/patients/animals/other materials for the study.

Presentation:

The information must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal.

9) References

The References must include only peer-reviewed sources (articles from academic journals and monographs) mentioned in the text of the article. It is not advised to references autoabstracts, dissertations, textbooks, study guides, standards, information from websites, statistic reports, publications in newspapers, websites and social media. If such information must be referred to, the source should be quoted in a footnote.

The description of a source should include its DOI, if it can be found (for foreign sources, that is possible in 95% of cases).

References to articles that have been accepted, but not yet published must be marked “in press”; the authors must obtain a written permission in order to reference such documents and confirmation that they have been accepted for publication. Information from unpublished sources must be marked “unpublished data/documents”; the authors also must obtain a written permission to use such materials.

References to journal articles must contain the year of publication, volume and issue, page numbers.

The description of each source must mention all of its authors.

The references, imprint must be verified according to the journals' or publishers' official websites.

Presentation:

References must be written in accordance with the Vancouver system.

The references must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with a 1.25-cm indent on the left. The font face must be normal. Each entry must be numbered in Arabic figures with a full stop after the number. The number with a full stop must be separated from the entry with a no-break space (Ctrl+Shift+Spacebar).

10) About the authors

Full second name, first name (in English); complete mailing address (including the postal code, city and country); complete name of the place of employment, position; academic degree, academic title, honorary degrees; membership in public associations, organizations, unions, etc.; official name of the organization in English; e-mail address; list and numbers of journals with the author's previous publications; the authors' photographs for publication in the journal.

Presentation:

The information must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal.

11) The authors' contribution

Detailed information as to each author's contribution to the article. For example: Author A analyzed literature on the topic of the paper, author B has developed a model of real-life facility operation, performed example calculation, etc. Even if the article has only one author, his/her contribution must be specified.

Presentation:

The information must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal.

12) Conflict of interests

A conflict of interests is a situation when people have conflicting and competing interests that may affect editorial decisions. Conflicts of interests may be potential or conscious, as well as actually existing. The objectivity may be affected by personal, political, financial, scientific or religious factors.

The author must notify the editorial office on an existing or a potential conflict of interests by including the corresponding information into the article.

If there is no conflict of interests, the author must also make it known. An example of wording: "The author declares the absence of a conflict of interests".

Presentation:

The text must be in 12-point Times New Roman, with a 1.5-line spacing, fully justified, with no indentation on the left. The font face must be normal.

SUBSCRIPTION TO THE JOURNAL «DEPENDABILITY»

Original article
Dependability, vol. 20 no. 2, 2020
<https://doi.org/10.21683/1729-2646-2020-20-2-43-53>

Application of machine learning methods for predicting hazardous failures of railway track assets

Igor B. Shubitskiy, Alexey M. Zamyshlav, Olga B. Pronovich, Alexey N. Ignatov, E. N. Ignatov
JSC MIAS, Russian Federation, Moscow, Moscow Aviation Institute, Russian Federation, Moscow, Russia
*shubitskiy@miast.com

Abstract. The aim of the paper is to reduce the number of hazardous failures of railway track assets by developing a method of prediction of rare hazardous failures based on the analysis of large amounts of data on each kilometre of track obtained in real time. Hazardous failures are rare events; the set of variable values of the number of hazardous failures on an individual kilometre of track per year is $\{0, 1\}$. However, for a railway track, the yearly number of such events is in the dozens and efficient analysis of the most probable location of failures. **Methods.** The problem of hazardous failures prediction is solved by means of machine learning. Hazardous events are predicted using the above statistics and artificial neural networks. Data Science technology is used. Such technology includes methods of data analysis, machine learning, and data mining. The application of various algorithms is demonstrated using the example of prediction of track superstructure failures. The result of facility rating is the conclusion regarding the location of hazardous failures of railway track. This conclusion is based on the comparison of the actual characteristics of an item and conditions of its operation with the characteristics of an item and conditions of its operation in the case of adverse events and cases of their non-occurrence. The practical value is the fact that the proposed set of methods and means can be used for the track maintenance decision-making system. It can be easily adapted and integrated into the automated measurement system installed on a railway track.

Keywords: machine learning; railway track facility failure; decision tree

For citation: Shubitskiy I.B., Zamyshlav A.M., Pronovich O.B., Ignatov A.N. Application of machine learning methods for predicting hazardous failures of railway track assets. Dependability, 2020, 20(2): 43-53. <https://doi.org/10.21683/1729-2646-2020-20-2-43-53>

Received on: 31.03.2020 / Upon revision: 18.04.2020 / For printing: 07.06.2020

As part of the initiatives aiming to promote the ideas of dependability in Russia and worldwide, the Editorial Board of the Dependability Journal will grant free access to any of its papers that you require.

Through the editorial office for any time-frame

tel.: +7 (495) 967-77-05;
e-mail: dependability@bk.ru

SUBSCRIBER APPLICATION FOR DEPENDABILITY JOURNAL

Please subscribe us for 20____
from No. _____ to No. _____ number of copies _____

Company name	
Name, job title of company head	
Phone/fax, e-mail of company head	
Mail address (address, postcode, country)	
Legal address (address, postcode, country)	
VAT	
Account	
Bank	
Account number	
S.W.I.F.T.	
Contact person: Name, job title	
Phone/fax, e-mail	

Publisher details: Dependability Journal Ltd.

Address of the editorial office: office 209, bldg 1, 27 Nizhegorodskaya Str., Moscow 109029,
Russia Phone/fax: 007 (495) 967-77-02, e-mail: dependability@bk.ru
VAT 7709868505 Account 890-0055-006
Account No. 40702810100430000017
Account No. 30101810100000000787

Address of delivery:

To whom: _____

Where: _____

To subscribe for Dependability journal, please fill in the application form and send it by fax or email.

In case of any questions related to subscription, please contact us.

Cost of year subscription is 4180 rubles, including 18 per cent VAT.

The journal is published four times a year.

**THE JOURNAL IS PUBLISHED WITH PARTICIPATION AND SUPPORT
OF JOINT-STOCK COMPANY RESEARCH & DESIGN INSTITUTE FOR INFORMATION
TECHNOLOGY, SIGNALLING AND TELECOMMUNICATIONS ON RAILWAY TRANSPORT
(JSC NIIAS)**



JSC NIIAS is RZD's leading company in the field of development of train control and safety systems, traffic management systems, GIS support technology, railway fleet and infrastructure monitoring systems



Mission:

transportation

□ efficiency,

□ safety,

□ reliability



Key areas of activity

- Intellectual control and management systems
- Transportation management systems and transport service technology
- Signalling and remote control systems
- Automated transportation management centers
- Railway transport information systems
- Geoinformation systems and satellite technology
- Transport safety systems
- Infrastructure management systems
- Power consumption and energy management systems
- Testing, certification and expert assessment
- Information security
- Regulatory support



www.vniias.ru