

INNOVATIVE TECHNOLOGIES FOR PROJECT AND PROGRAM MANAGEMENT



EIROPAS SAVIENĪBA

Riga – 2025



INNOVATIVE TECHNOLOGIES FOR PROJECT AND PROGRAM MANAGEMENT

*Collective monograph
edited by I. Linde*

European University Press
ISMA University of Applied Sciences
Riga (Latvia) 2025

LĒMUMU ATBALSTA SISTĒMAS PROJEKTU UN PROGRAMMU VADĪBĀ

Kolektīvas monogrāfija
I. Linde zinātniskajā redakcijā

Eiropas Universitāte Press
Informācijas sistēmu menedžmenta augstskola
Rīga (Latvija) 2025

DOI: <https://doi.org/10.66118/MMP.2025>

ISBN 978-9984-891-34-7

Innovative Technologies for Project and Program Management [Text]: Collective monograph edited by I. Linde. European University Press. Riga: ISMA, 2025. 275 p.

Reviewers:

Oleg Fedorovich – Dr. Sc. (Engineering), Professor, Head of Department of Computer Sciences and Information Technologies, National Aerospace University «Kharkiv Aviation Institute».

Heorhii Kuchuk – Dr. Sc. (Engineering), Professor, Professor of the Department of Computer Engineering and Programming, National Technical University «Kharkiv Polytechnic Institute».

Authors: Anishchenko A., Artomova A., Baryshevskyi A., Bazilevych K., Bulavin D., Cherkasov D., Chumachenko D., Chumachenko T., Danshyna S., Dobrytskyi D., Farionova T., Fedorovich O., Fedorovych V., Fonarova T., Gubka O., Khrustalova S., Kobylkin D., Kosenko N., Kosenko V., Kovalchuk O., Kozyr S., Kritskiy D., Kuznetsova Yu., Leshenko Yu., Malieiev L., Malyeyeva O., Menailov Ie., Moiseienko A., Molokanova V., Neronov S., Oriekhov O., Parfeniuk Yu., Petrenko V., Plehov D., Plehova G., Podorozhko K., Popov A., Prohorov O., Sharonova N., Sopov Ie., Timofeyev V., Troshchy D., Vyprytskyi A.

The monograph presents the achievements of Ukrainian scientists in the field of business management, use of economic and mathematical modeling, information technologies, management technologies and technical means in the field of functioning, development, and project management at enterprises.

The publication is recommended for professionals in the fields of economics, information technology, project and program management – for undergraduate and graduate students, as well as academics and teachers of higher education.

The articles are reproduced from the original authors, in the author's edition.

DOI: <https://doi.org/10.66118/MMP.2025>

ISBN 978-9984-891-34-7

CONTENT

7	INTRODUCTION
8	Adaptive control under uncertainty <i>Anishchenko A., Khrustalova S., Timofeyev V.</i>
30	Digital transformation project management: innovative approaches and tools <i>Baryshevskyi A.</i>
37	Integration capabilities for enterprise value chain management <i>Bulavin D., Petrenko V.</i>
47	Experimental evaluation of modern UI/UX design tools as a component of project management <i>Cherkasov D., Kuznetsova Yu.</i>
71	Integrating Epidemic Simulation into Crisis Management Programs: Lessons from Ukraine <i>Chumachenko D., Meniailov Ie., Bazilevych K., Parfeniuk Yu., Chumachenko T.</i>
84	Infrastructure projects: assessment of anthropogenic impact on river ecosystems <i>Danshyna S., Podorozhko K.</i>
94	Multi-Criteria Decision-Making Approaches for IT Outsourcing Project Portfolio Management <i>Dobrytskyi D., Fonarova T.</i>
103 ...	Modeling the formation of drone swarms using component and agent-based approaches <i>Fedorovich O., Prohorov O., Leshenko Yu., Malieiev L. Kosenko V.</i>
134 ...	Digitization of IT team management processes in project management: the role of information systems and artificial intelligence <i>Kovalchuk O., Kobylkin D.</i>
143 ...	Management of monitoring and control processes in innovative projects of uav parameterization and swarm flight: models, strategies, adaptive solutions <i>Kritskiy D., Malyeyeva O., Popov A., Gubka O., Kosenko N.</i>

- 171 ...Deployment Strategies and Architectural Innovations
in Learning Management Systems
Moiseienko A., Kuznetsova Yu.
- 183 ...Methods of introducing digital transformation
into the organizations development
Molokanova V. Kozyr S.
- 192 ...Information technology for estimating the size of software projects
Oriekhov O., Farionova T.
- 202 ...Mathematical bases of the methodology for building an intellectual system
based on the theory of intelligence and the apparatus of predicate algebra
Plehova G., Sharonova N., Neronov S., Plehov D., Fedorovych V.
- 235 ...Deep reinforcement learning-based robust routing
for scalable UAV swarm communication
Sopov Ie., Artomova A.
- 248 ...Best practices for systematizing routine processes
in mobile development projects
Troshchyi D., Kuznetsova Yu.
- 258 ...Coercive measures applied under the legal emergency regime
Vypriyskyi A.

INTRODUCTION

The key to the successful operation of complex socio-economic and technical systems is their constant updating, adaptation to changing environmental conditions and appropriate self-regulation of the internal structure, processes and technologies.

The scientific and methodological developments proposed in the monograph relate to various areas of strategic development, the use of modeling and information technologies, project and program management technologies. They will help to improve existing processes and develop new ones in various sectors of the economy and production.

The issues of applying information technology in the management of dynamic objects in various industries are considered, each of which is characterized by the allocation of objects, management methods and their impact on the project as a way of organizing work to achieve a goal or solve a problem. Typical tasks that can be identified in project management are considered. Much attention is paid to the use of artificial intelligence methods.

ADAPTIVE CONTROL UNDER UNCERTAINTY

Anishchenko A., Khrustalova S., Timofeyev V.

Rapid progress of computing technology and significant increase of microprocessor performance open new opportunities for real-time control of complex objects. First of all, we are talking about adaptive control systems in the broad sense, which use in real time the incoming information about signals and parameters of the object under study to develop control actions under conditions of significant uncertainty about the object and disturbances acting on it.

One of the fundamental problems of control theory is the problem of control, i.e., determining or estimating the system parameters at different points in time. Knowledge of system parameters plays an important role in system theory when solving control, observation, forecasting and diagnostic problems [2].

The task of identification and the closely related task of control and detection of changes in the properties of an object consists in the construction of an optimal in some sense model of the system (object) based on the results of observations of input and output coordinates of the system, i.e., a formalized representation of this system [3, 4,] and the detection of changes (control) using this model.

In practice, the so-called identification problem in a narrow sense often arises. The identification problem in this sense consists in estimating the parameters and state of the system based on the results of observations of input and output variables obtained under conditions of normal functioning of the object. In this case, the structure of the system is known and the class of models to which the object belongs is specified.

A priori information about the object and its operating conditions is often absent or very poor, so additional tasks have to be solved beforehand. These tasks include: studying the nature and character of external influences, estimation of noise boundaries, early technical diagnostics and control of the object. In this case, control implies recognizing the state of the technical system and detecting changes in the properties of the functioning object [5].

The proposed work solves the problem of identification, control and management of dynamic stochastic systems under conditions of normal operation and significant a priori uncertainty about the noise affecting the object (there is a complete lack of information about the nature and character of the noise) on the basis of an unconventional adaptive approach.

1. Mathematical models of dynamic stochastic objects

Dynamic objects are described by differential, integral and functional equations with respect to some coordinates characterizing their state [5]. The use of computational tools to control such objects leads to the necessity to describe the studied objects by differences or differential equations.

Linear dynamic stochastic objects (Fig. 1) in general case can be represented by a linear difference equation:

$$y(k) + \sum_{i=1}^n a_i y(k-i) = \sum_{i=0}^m b_i u(k-i) + \sum_{i=0}^l c_i w(k-i), \quad (1)$$

where $y(k)$ – output signal, $u(k)$ – input signal, $w(k)$ – perturbation (disturbance); $a_i (i=1, \dots, n)$, $b_i (i=0, 1, \dots, m)$, $c_i (i=0, 1, \dots, l)$ – object parameters to be determined; $k=0, 1, 2, \dots$ – discrete time.

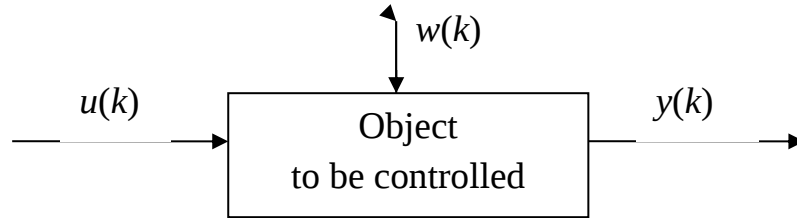


Fig. 1. Linear dynamic stochastic object in general form

Introduction of the delay operator q^{-1} , defined in accordance with the condition

$$q^{-d} y(k) = y(k-d), \quad (2)$$

allows to represent equation (1) in the operator form

$$A(q^{-1})y(k) = B(q^{-1})u(k) + C(q^{-1})w(k) \quad (3)$$

or

$$y(k) = K_u(q^{-1})u(k) + K_w(q^{-1})w(k), \quad (4)$$

where

$$K_u(q^{-1}) = B(q^{-1})/A(q^{-1}), \quad K_w(q^{-1}) = C(q^{-1})/A(q^{-1}) \quad (5)$$

are the transfer functions of the object under control and perturbation, respectively.

If the dynamic object is stable (and below this assumption is used everywhere), then the roots of the characteristic polynomial are as follows

$$A(q^{-1}) = 1 + a_1 q^{-1} + a_2 q^{-2} + \dots + a_n q^{-n}$$

lie inside the unit circle centered at the origin.

Moreover, the dynamic object (1) is minimal-phase if the roots of polynomials

$$\begin{aligned} B(q^{-1}) &= b_0 + b_1 q^{-1} + b_2 q^{-2} + \dots + b_m q^{-m}, \\ C(q^{-1}) &= c_0 + c_1 q^{-1} + c_2 q^{-2} + \dots + c_l q^{-l} \end{aligned}$$

lie inside the unit circle centered at the origin, and non-minimal-phase otherwise.

The task of structural identification consists in determining the order of the polynomials used in the description. The parametric identification task consists in determining the parameters of polynomials whose order is already known or estimated. In this case, all parameters to be determined, are included in some vector of unknown parameters.

$$y(k) = K_u(q^{-1}, \Theta)u(k) + K_w(q^{-1}, \Theta)w(k). \quad (6)$$

Equations (1), (3), (4) are the most general and cover as special cases various objects and processes (regression, autoregression, moving average, etc.).

Let us focus on some common ways of describing dynamic objects. Let's consider an autoregressive model to describe dynamic objects. Autoregressive models or ARX models have been investigated in [6]. ARX models are described in the form of the following difference equation

$$y(k) + a_1 y(k-1) + \dots + a_n y(k-n) = b_1 u(k) + \dots + b_m u(k-m) + w(k) \quad (7)$$

or

$$A(q^{-1})y(k) = B(q^{-1})u(k) + w(k). \quad (8)$$

In equation (5), white noise enters as an unobservable quantity. If we denote «and», we can see that (8) coincides with (6).

To represent a nonstationary dynamic object by an ARX model, the following relation is used:

$$\begin{aligned} y(k) + a_1(k)y(k-1) + \dots + a_n(k)y(k-n) &= \\ = b_1(k)u(k-1) + \dots + b_m(k)u(k-m) + w(k) \end{aligned} \quad (9)$$

or

$$A(q^{-1}, \Theta(k))y(k) = B(q^{-1}, \Theta(k))u(k) + w(k). \quad (10)$$

This model contains a time-dependent vector of adjustable parameters of the form

$$\Theta(k) = (a_1(k), \dots, a_n(k), b_1(k), \dots, b_m(k))^T. \quad (11)$$

By labeling

$$\varphi(k) = (-y(k-1), \dots, -y(k-n), u(k-1), \dots, u(k-m))^T, \quad (12)$$

we can rewrite (10) as follows:

$$y(k) = \Theta^T(k) \varphi(k) + w(k). \quad (13)$$

The obtained relation (13) describes a model of pseudo-linear regression type, for the study of which it is possible to apply developed methods of linear regression analysis [6]. It should be noted that the description of a linear dynamic object in the form of (13) is currently widespread, as it allows to represent the identification and control algorithm in a fairly simple and visual form.

Along with autoregressive models, autoregressive moving average models (ARMAX-models) are widely used. ARMAX-models were introduced to solve system identification problems by Ostrem and Bolin [7] and since then such models are among the most popular structures. Various modifications of ARMAX-models are given in [8].

The main difference between the ARMAX model and the ARX model (9) is to give more flexibility in describing the properties of the disturbance. In this model, the model error is represented as a moving average of white noise, which leads to the following relation:

$$\begin{aligned} y(k) + a_1(k)y(k-1) + \dots + a_n(k)y(k-n) = \\ = b_1(k)u(k-1) + \dots + b_m(k)u(k-m) + \\ + w(k) + c_1(k)w(k-1) + \dots + c_l(k)w(k-l). \end{aligned} \quad (14)$$

If we denote

$$C(k, q^{-1}) = 1 + c_1(k)q^{-1} + \dots + c_l(k)q^{-l},$$

then equation (14) can be written in the form of

$$A(k, q^{-1})y(k) = B(k, q^{-1})u(k) + C(k, q^{-1})w(k), \quad (15)$$

which leads to (6) if we put

$$\begin{aligned} K_u(k, q^{-1}, \Theta) &= B(k, q^{-1}, \Theta) / A(k, q^{-1}, \Theta), \\ K_w(k, q^{-1}, \Theta) &= C(k, q^{-1}, \Theta) / A(k, q^{-1}, \Theta) \end{aligned}$$

Then

$$\Theta(k) = (a_1(k), \dots, a_n(k), b_1(k), \dots, b_m(k), c_1(k), \dots, c_l(k))^T.$$

In this case, initialization of the estimation procedure at the moment $k = 0$ requires knowledge of the following quantities $y(0), \dots, y(-k^* + 1)$, $k^* = \max(n, l)$, $\hat{y}(0), \dots, \hat{y}(-l + 1)$, where $\hat{y}(k)$ is the estimate of the model output signal.

If these quantities are missing, they can be taken equal to zero, which introduces inaccuracy into the model. It is also possible to start the recursion at the moment $\max(k^*, u(l))$ and include the unknown initial conditions $y(k), k=1, \dots, l$ in the vector Θ .

Having introduced the notations

$$\varepsilon(k) = y(k) - \hat{y}(k) \quad (16)$$

$$y(k) = (-y(k-1), \dots, -y(k-n), u(k-1), \dots, u(k-m), \varepsilon(k-1), \dots, \varepsilon(k-l))^T \quad (17)$$

expression (15) can be written as follows:

$$y(k) = \Theta^T(k) \varphi(k) + w(k). \quad (18)$$

The variable used in the relationship $\varepsilon(k)$ is called the prediction error or generalized residual. Equation (18) itself is called a pseudolinear regression $\varphi(k)$ due to the nonlinear dependence of the vector on $\Theta(k)$.

The structures discussed above can actually lead to a large number of models, depending on which of the sets A, B, C are included in the model. For convenience, a generalized model structure for describing a non-stationary object can be used:

$$A(k, q^{-1})y(k) = \frac{B(k, q^{-1})}{\tilde{B}(k, q^{-1})}u(k) + \frac{C(k, q^{-1})}{\tilde{C}(k, q^{-1})}w(k)$$

where

$$\begin{aligned} \tilde{B}(k, q^{-1}) &= 1 + \tilde{b}_1(k)q^{-1} + \dots + \tilde{b}_l(k)q^{-l}; \\ \tilde{C}(k, q^{-1}) &= 1 + \tilde{c}_1(k)q^{-1} + \dots + \tilde{c}_l(k)q^{-l}. \end{aligned}$$

The variable $\varepsilon(k)$ used in the relationship is called the prediction error or generalized residual. Equation (18) itself is called a pseudolinear regression due to the nonlinear dependence of the vector $\varphi(k)$ on $\Theta(k)$.

The structures discussed above can actually lead to a large number of models, depending on which of the sets are included in the model. A, B, C .

For convenience, a generalized model structure for describing a non-stationary object can be used:

$$A(k, q^{-1})y(k) = \frac{B(k, q^{-1})}{\tilde{B}(k, q^{-1})}u(k) + \frac{C(k, q^{-1})}{\tilde{C}(k, q^{-1})}w(k), \quad (19)$$

where

$$\begin{aligned}\tilde{B}(k, q^{-1}) &= 1 + \tilde{b}_1(k)q^{-1} + \dots + \tilde{b}_l(k)q^{-l}; \\ \tilde{C}(k, q^{-1}) &= 1 + \tilde{c}_1(k)q^{-1} + \dots + \tilde{c}_l(k)q^{-l}.\end{aligned}$$

However, for most practical purposes this structure is too general. In applications, one or more of the five polynomials are usually assumed to be equal to 1. Using the above $\Theta(k)$ vector of generally non-stationary parameters and the generalized history vector $\varphi(k)$ allows us to represent model (19) as a pseudo-linear regression (18). If the input signals $u(k)$ and outputs $y(k)$ are vectors $(p_u \times 1)$ and $(p_y \times 1)$ respectively, then the simplest generalization of the set of models (6) is

$$\begin{aligned}y(k) + A_1(k)y(k-1) + \dots + A_n(k)y(k-n) = \\ = B_1(k)u(k-1) + \dots + B_m(k)u(k-m) + w(k),\end{aligned}\quad (20)$$

where $A_i(k), B_i(k)$ are the matrices of dimensions $p_y \times p_y$ and $p_y \times p_u$, respectively. Similarly, matrix polynomials of the variable are introduced: of the models (6) there are

$$\begin{aligned}A(k, q^{-1}) &= A_1(k)q^{-1} + \dots + A_n(k)q^{-n}; \\ B(k, q^{-1}) &= B_1(k)q^{-1} + \dots + B_m(k)q^{-m}.\end{aligned}$$

In this case, we get not a vector, but a $\Theta(k)$ matrix of the required parameters of dimension $((np_y + mp_u)p_y)$:

$$\Theta(k) = [A_1(k), A_2(k), \dots, A_n(k), B_1(k), \dots, B_m(k)]^T \quad (21)$$

and

$$\varphi(k) = [-y^T(k-1), \dots, -y^T(k-n), u^T(k-1), \dots, u^T(k-m)]^T - \quad (22)$$

a column vector of dimension $(n \times p_y + m \times p_u) \times 1$, which allows us to represent (21) by the matrix equation of linear regression (13).

These formulas can be considered as a set of linear regressions written one under the other with the same input vector.

Thus, as follows from the above, the description of a p_y linear dynamic stochastic object by equation (13) is quite general. On the other hand, the representation of the object under study by model (13) allows the use of

well-developed methods of regression analysis and methods for identifying static objects for identifying parameters.

These formulas can be viewed as a set of linear regressions written one under the other with the same input vector. However, the application of regression analysis methods requires certain assumptions regarding the statistical properties of useful signals and interference. In particular, it is assumed that the interference is a random Ψ vector with zero mathematical expectation and a diagonal covariance matrix.

$$M\{\Psi(k)\} = 0, \quad i = 1, 2, \dots, k, \quad \text{cov}\{\Psi(k)\} = \sigma_w^2 I. \quad (23)$$

Here $M\{\bullet\}$ – the symbol is the symbol of mathematical expectation,

σ_w^2 – is the variance of noise,

I – is the identity matrix of dimension $k \times k$.

All limitations associated with classical regression analysis are determined by the traditional criterion of squared errors and, accordingly, the least squares method, which is quite rigidly tied to the normal distribution law of variables and disturbances. If, however, it is only known about the interference that it is limited in amplitude

$$|w(k)| < \delta, \quad (24)$$

then other methods are used to estimate the parameters, which do not use any information about the statistical properties of disturbances except for their belonging to a certain limited interval.

Currently, for estimating parameters in the presence of limited-amplitude interference, the most widely used algorithms are those based on the polytope method (and, as a special case, orthotopes), algorithms containing a dead zone, and algorithms based on constructing ellipsoids. This direction is currently developing rapidly and attracts much research.

2. Recursive identification algorithms

We consider an object described by the pseudolinear regression equation (13):

$$y(k) = \Theta^T(k) \varphi(k) + w(k)$$

with noise satisfying constraint (24)

$$|w(k)| < \delta.$$

Then equation (13) can be rewritten as

$$y(k) = \Theta^T(k) \varphi(k) + w(k)$$

with a vector of sought parameters satisfying the system of inequalities

$$|w(k)| < \delta.$$

Then equation (13) can be rewritten as

$$\varepsilon(k) = y(k) - \Theta^T(k)\varphi(k)$$

with a vector $\hat{\Theta}(k)$ sought parameters satisfying the system of inequalities

$$\left(y(k) - \Theta^T(k)\varphi(k)\right)^2 \leq \delta^2 \quad \forall k = 1, 2, \dots$$

Only those that belong to the set can be used as estimates.

With the vector of the sought parameters satisfying the system of inequalities.

Only those that belong to the set can be used as estimates:

$$S(k) = \left\{ \Theta : y(k) - \Theta^T \varphi(k) \leq \delta \right\}. \quad (25)$$

The undefined parameters belong to the domain bounded by the hyperplanes (25). The domain is a simple connected domain. The space is a convex polytope, which can become quite complex if n is large (Fig. 2).

From a geometric point of view, the set $S(k)$ is a monotonically increasing sequence of intersections of convex polytopes $F(k)$

$$S(k) = \bigcap_{j=1}^k F_j = S(k-1) \cap F(k) \quad (26)$$

$$F(k) = \left\{ \Theta : |y(k) - \Theta \varphi(k)| \leq \delta \right\}. \quad (27)$$

Calculating estimates $\hat{\Theta} \in S(k)$ is a complex problem, the solution of which can be significantly simplified by constructing a certain set that restricts $S(k)$ [10]. These restrictions can be specified in the form of orthotopes [11]. The problem of constructing an identification algorithm consists of finding a procedure for recurrent refinement based on the existing values when new observations of input and output signals arrive $S(k)$. The region of uncertain parameters, limited by hyperplanes (25), can be approximated by an orthotope whose edges are parallel to the axes. A geometric illustration of this approach is shown in Fig. 3.

The boundaries $\Theta_{\min}$, $\Theta_{\max}$ of the indefinite interval are associated with the i -th parameter and give an idea of the minimum and maximum values of the criterion

$$J_i(\Theta) = \Theta_i, \quad (28)$$

where A is the permissible area of parameters Θ is the area $S(k)$.

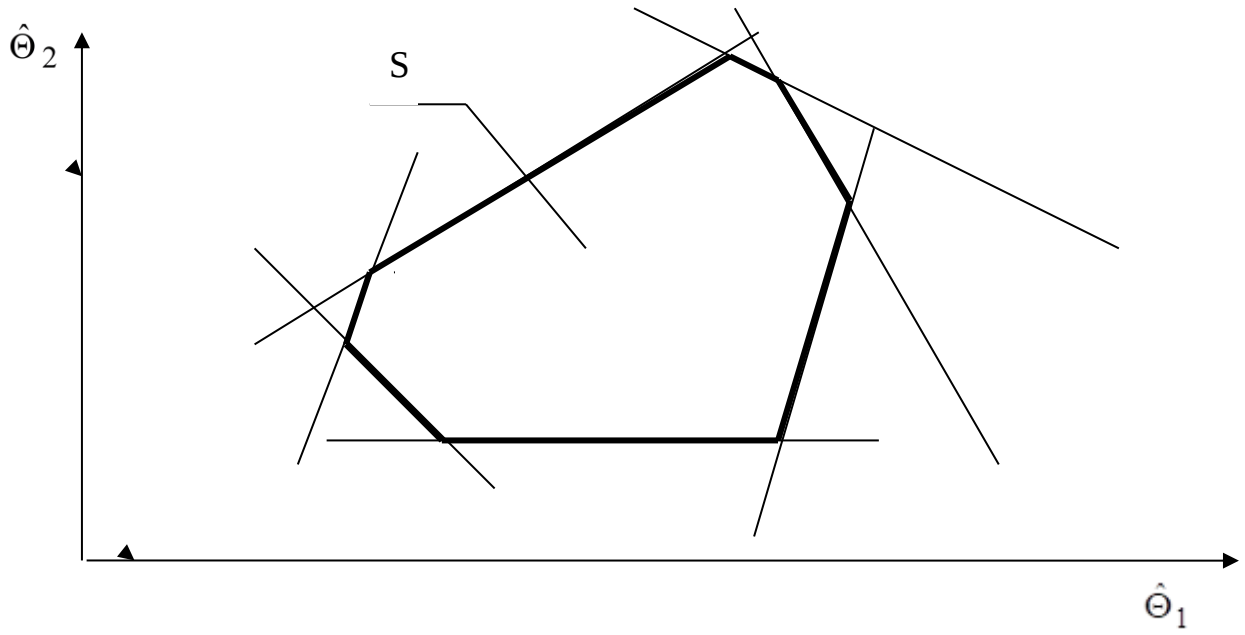


Fig. 2. Finding the range of acceptable values

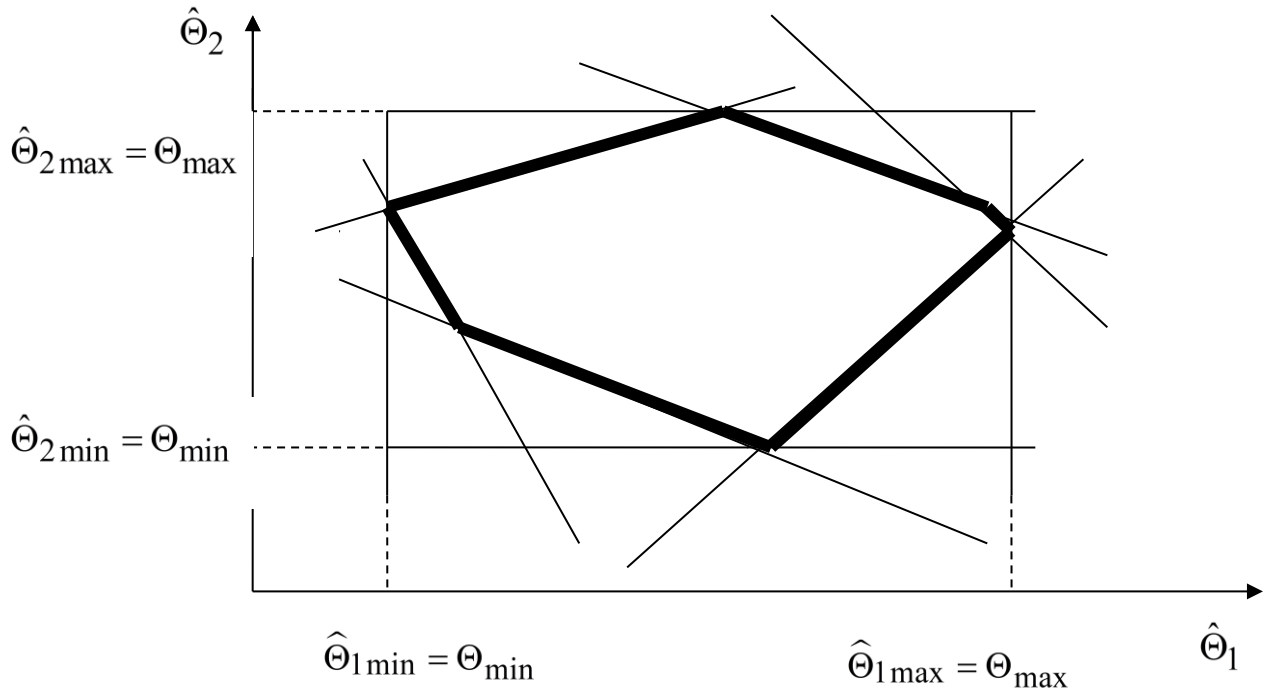


Fig. 3. Approximation of an indefinite interval by an orthotop

Criterion (28) is determined by a set of linear inequalities. Each boundary, therefore, is obtained by solving the problem of linear programming [9, 10]. Calculating the boundaries of the orthotope requires the solution of linear programming tasks, where – the size of the space. Each of the linear programming problems contains linear inequalities. The resulting algorithm is not recurrent

and involves a large number of more complex calculations than algorithms of the method of least squares (MHK).

When the area $S(k)$ is narrow and oriented at a large angle to the axes, limiting the axial orthotope can give very inaccurate results. However, it is possible to try to work with an orthotope of minimal volume.).

If it is known about the interference that they satisfy the condition (24), then this information can be taken into account in the algorithm containing the zone of insensitivity. At the same time, the requirements regarding the knowledge of the properties of interference are reduced, but the identification algorithm itself is coarsened.

The idea of using the insensitivity zone in the algorithm is based on the fact that even in the case of accurate determination of the model parameters, there remains an error (mismatch between the output signals of the object and the model), the magnitude of which is determined by the magnitude of the noise (24). On the identification algorithm it will be as follows. The magnitude of the error $\Theta(k-1) = \Theta(k)$ in the case is equal

$$y(k) - \Theta^T(k) \varphi(k) = w(k),$$

as follows. The error value in the case of $\Theta(k-1) = \Theta(k)$ is equal to

$$y(k) - \Theta^T(k) \varphi(k) = w(k),$$

that is, on the $(k-1)$ -m step, the correction of the estimate $\Theta(k-1)$ will be made on the value, defined $w(k)$ and on the $(k-1)$ -om step, the $\Theta(k-1)$ estimate will be obtained. To prevent this from happening, a zone of insensitivity is introduced, in which case

$$\left| y(k) - \Theta^T(k) \varphi(k) \right| \leq \delta(k)$$

that is, on the second step, the correction of the estimate will be made on the value determined and on the second step, the estimate will be obtained. To prevent this from happening, a zone of insensitivity is introduced, in which case

$$\xi(\varphi) = \begin{cases} \varphi - \delta & \text{at } \varphi > \delta; \\ 0 & \text{at } |\varphi| \leq \delta; \\ \varphi + \delta & \text{at } \varphi < -\delta, \end{cases}$$

no assessment correction occurs when.

Various algorithms including deadband type nonlinearity were considered in the works [11–13].

If the noise distribution density is strictly limited in the interval $(-\delta, \delta)$, then these algorithms converge. The estimation procedure itself is reduced to the sequential projection of current estimates $\Theta(k-1)$ into a system of infinite bands limited by parallel planes and determined by inequalities in the parameter space

$$-\delta \leq y(k) - \Theta^T(k-1)\varphi(k) \leq \delta.$$

Thus, in [14] the following algorithm was considered:

$$\Theta(k) = \Theta(k-1) + \frac{\xi\left(y(k) - \Theta^T(k-1)\varphi(k)\right)}{\|\varphi(k)\|^2} \varphi(k), \quad (29)$$

which can be rewritten as: which can be rewritten a

$$\Theta(k) = \Theta(k-1) + \rho \frac{\text{sign}\left|y(k) - \Theta^T(k-1)\varphi(k)\right|}{\|\varphi(k)\|^2} \varphi(k), \quad (30)$$

where is $\rho(k) = \left|\xi\left(y(k) - \Theta^T(k-1)\varphi(k)\right)\right|$ is the distance from the point $\Theta(k-1)$ to the strip:

$$-r \leq y(k) - \Theta^T(k-1)\varphi(k) \leq r$$

where $\rho(k) = \left|\xi\left(y(k) - \Theta^T(k-1)\varphi(k)\right)\right|$ – distance from the point $\Theta(k-1)$ to the strip $-r \leq y(k) - \Theta^T(k-1)\varphi(k) \leq r$.

In [15] the dynamics of the Strip algorithm was studied:

$$\Theta(k) = \Theta(k-1) + \gamma w(k) \frac{y(k) - \Theta^T(k)\varphi(k)}{\|\varphi(k)\|^2} \varphi(k), \quad (31)$$

where

$$w(k) = \begin{cases} 0, & \text{if } \left|y(k) - \Theta^T(k-1)\varphi(k)\right| \geq \delta; \\ 1, & \text{if } \left|y(k) - \Theta^T(k-1)\varphi(k)\right| < \delta. \end{cases} \quad (32)$$

In the work [16] the problem of constructing an algorithm with a dead zone that minimizes the criterion was considered

$$J(\Theta(k)) = \frac{1}{2} \|\Theta(k) - \Theta(k-1)\|^2 \quad (33)$$

under restriction

$$|e(k)| \leq \delta(k) \quad (34)$$

and provided that the a posteriori error $e^0(k)$ satisfies the relation

$$(e^0(k) - i^0(k)r(k)i(k)) = 0. \quad (35)$$

Here

$$e^0(k) = y(k) - \Theta^T(k)\varphi(k), \quad (36)$$

$$i(k) = \begin{cases} 0, & \text{if } |e(k)| \leq \delta; \\ 1, & \text{if } |e(k)| > \delta, \end{cases} \quad (37)$$

$$i^0(k) = \text{sign}(e(k)). \quad (38)$$

As can be seen from (36), the difference between the a posteriori error $e^0(k)$ and $e(k)$ is that when calculating $e^0(k)$ at the k -th step, the estimate $\Theta(k)$ obtained at the k same step is used, and not $\Theta(k-1)$, used in $e(k)$.

In this case, the problem of constructing an algorithm is reduced to minimizing the Lagrange function of the form:

$$J(\Theta(k), \lambda(k)) = \frac{1}{2} \|\Theta(k) - \Theta(k-1)\|^2 + \lambda(k)(e^0(k) - e^0(k)\delta(k))i(k), \quad (39)$$

where is the Lagrange multiplier $\lambda(k)$.

From the minimum condition it follows

$$\frac{\partial(\Theta(k)\lambda(k))}{\partial\Theta} = 0; \quad \frac{\partial(\Theta(k)\lambda(k))}{\partial\lambda} = 0, \quad (40)$$

from which we obtain a system of equations, solving which we find

$$\Theta(k) = \Theta(k-1) + \lambda(k)\varphi(k)i(k), \quad (41)$$

$$\lambda(k) = \begin{cases} (e(k) - i(k)r(k)) / \|\varphi(k)\|^2, & \text{if } \varphi(k) \neq 0; \\ 0, & \text{if } \varphi(k) = 0. \end{cases} \quad (42)$$

The procedure (41)–(42) is a variation of the Kaczmarz algorithm, implementing a combination of the normalized least squares algorithm and the projection algorithm with a dead zone [18, 19].

The error functions used in the algorithms, taking into account the dead zone, are shown in Fig. 4, 5.

The dead zone shown in Fig. 4 is described by the relations.

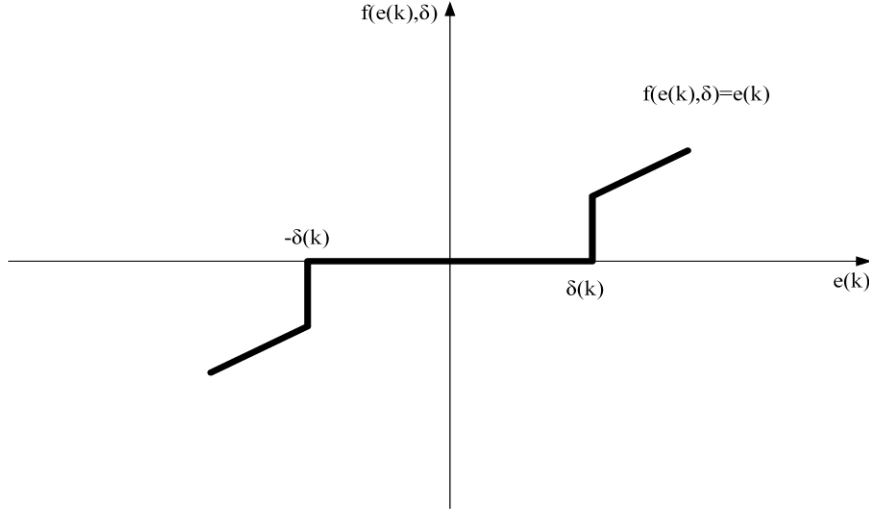


Fig. 4. The dead zone described by (42)

The use of elements implementing the dead zone shown in Fig. 5 and described by the relation has become widespread.

$$f(e(k), \delta) = \begin{cases} e(k) - \delta(k) \text{signe}(k), & \text{if } |e(k)| \geq \delta; \\ 0, & \text{if } |e(k)| < \delta. \end{cases} \quad (44)$$

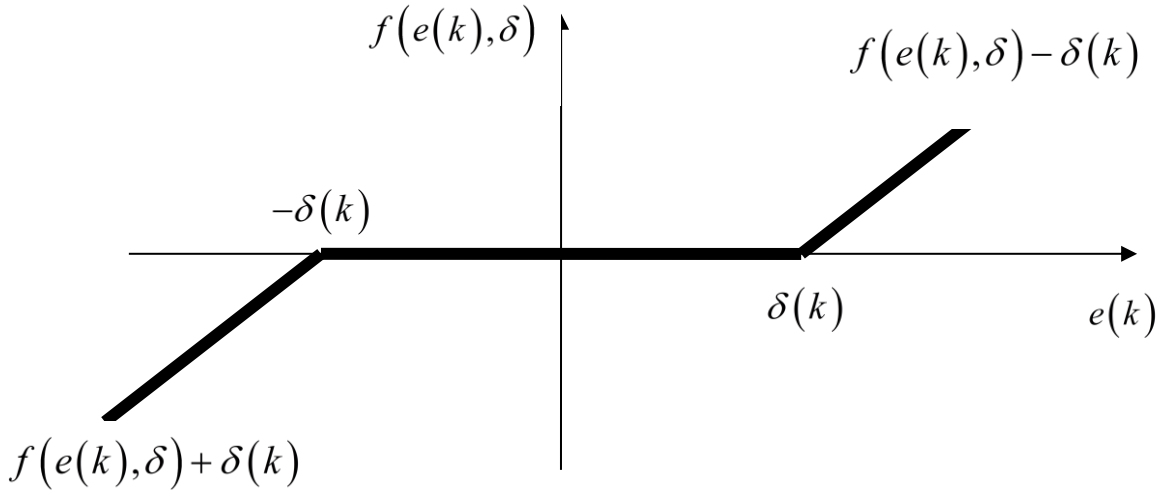


Fig. 5. The dead zone described by (44)

As noted earlier, the calculation of estimates $\Theta \in S(k)$ is a complex problem, the solution of which can be significantly simplified by constructing a certain set, $S(k)$ bounding in the form of ellipsoids [20], the center of which coincides with $\Theta(k)$. In identification problems, algorithms based on the construction of ellipsoids and using the following properties of a unit ellipsoid have become widespread [21]:

$$E = \left\{ x \in R^p : x^T A^{-1} x \leq 1 \right\}, A = A^T > 0$$

- the axes of the ellipsoid form an orthogonal system of eigenvectors of the A matrix, and the $\lambda_i^{1/2}$ semi-axes are equal to, where are the λ_i eigenvalues $A(i=1,2,...,p)$;

- the volume of the ellipsoid is proportional to the product of the semi-axes, i.e. proportional to $\sqrt{\det A}$;

- the sum of the squares of the semi-axes $\sum_{i=1}^k \left(\lambda_i^{1/2}\right)^2$ is equal to the trace of the matrix A .

Let

$$E(k-1) = \left\{ \Theta : \left(\Theta - \Theta(k-1) \right)^T P^{-1}(k-1) \left(\Theta - \Theta(k-1) \right) \leq r^2 \right\} \quad (45)$$

some ellipsoid that contains the intersection, i.e. $S(k-1)$. Here $S(k-1) \in E(k-1)$ is the center of the ellipsoid $S(k-1) \in E(k-1)$ and is the matrix defining the semi-axes $E(k-1)$ (i.e. the orientation and dimensions of the ellipsoid). The matrix $P(k-1)$ is symmetrically permutable, and $P(k-1) = P^T(k-1) > 0$. For degenerate ellipsoids, the constraints correspond to two parallel hyperplanes, $P(k-1) = P^T(k-1) > 0$ where is determined non-uniquely and is a matrix $P(k-1)$ with unit rank.

The center of the ellipsoid $\Theta(k-1)$ is an estimate of the desired parameter vector Θ at time $k-1$. The arrival of new measurements and gives the following value, and the intersection (26) due to the fact that $S(k-1) \in E(k-1)$, is completely included in the intersection $E(k-1) \cap F(k)$. The new ellipsoid is constructed so that.

$$\{E(k-1) \cap F(k)\} \subset E(k). \quad (46)$$

In addition, the new assessment must provide the condition

$$E_k \leq E_{k-1}. \quad (47)$$

Thus, the task of constructing an algorithm consists of finding such a recurrent refinement based $\Theta(k), P^{-1}(k), r^2$ on the existing values upon $\Theta(k), P^{-1}(k), r^2$ receipt of new observations of input and output signals that (46) and (47) are satisfied. A geometric illustration of this approach is shown in Fig. 6.

The ellipsoid $E(k)$ is constructed with initial conditions $E(k)$, where I is the identity matrix of size $p \times p$, and is α some sufficiently large number. This means

that the initial ellipsoid E_0 is a sphere of radius $\alpha^{1/2}$. The expression for the construction $E(k)$ is as follows [22]:

$$\begin{aligned} \Theta_0 &= \Theta, \quad P_0 = \alpha I, \\ \Theta(k) &= \Theta(k-1) + \frac{P(k-1)(y(k) - \varphi^T(k)\Theta(k-1))}{1 + \varphi^T(k)P(k-1)\varphi(k)}\varphi(k), \end{aligned} \quad (48)$$

$$P(k) = P(k-1) - \frac{P(k-1)\varphi(k)\varphi^T(k)P(k-1)}{1 + \varphi^T(k)P(k-1)\varphi(k)}, \quad (49)$$

which is identical in structure to the recurrent least squares method (RLSM).

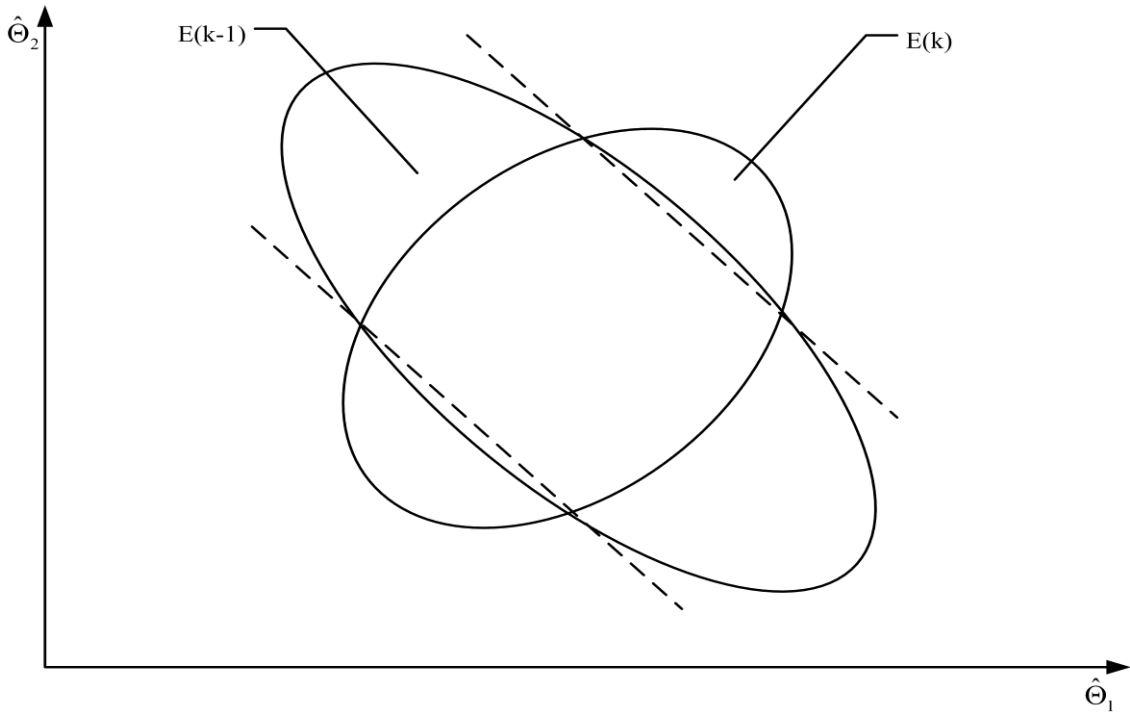


Fig. 6. Ellipsoid constraint

To minimize the geometric dimensions of an ellipsoid $E(k)$, criteria such as minimum volume or minimum trace are used. The resulting algorithms are called, respectively, the «minimum volume algorithm» (or «volume-minimizing algorithm») and the «minimum trace algorithm» (or «trace-minimizing algorithm»).

In [22], a recurrent procedure for calculating an ellipsoid $E(k)$ was proposed

In the work [22], a recurrent procedure for calculating the ellipsoid was proposed. To minimize the geometric dimensions of an ellipsoid, criteria such as minimum volume or minimum trace are used. The resulting algorithms are called, respectively, the «minimum volume algorithm» (or «volume-minimizing algorithm») and the «minimum trace algorithm» (or «trace-minimizing algorithm»).

$$\begin{aligned}
E(k) = & \left\{ \Theta : \alpha(k) \left(\Theta - \Theta(k-1) \right)^T P^{-1}(k-1) \left(\Theta - \Theta(k-1) \right) + \right. \\
& \left. + \beta(k) \left(y(k) - \Theta^T(k-1) \varphi(k) \right)^2 \leq \alpha^2(k) + \beta(k) r^2 \right\} = \\
= & \left\{ \Theta : \left(\Theta - \Theta(k-1) \right)^T P^{-1}(k-1) \left(\Theta - \Theta(k-1) \right) \leq r^2 \right\},
\end{aligned} \tag{50}$$

where $\alpha(k) \in (0,1)$ is the parameter for forgetting outdated information; is the parameter for weighing (selecting) newly incoming information

where $\alpha(k) \in (0,1)$ – obsolescence parameter;

$\beta(k) \in (0,1)$ – weighting (selection) parameter for newly received information.

The task of constructing an algorithm is reduced to finding the center $\Theta(k)$ of an ellipsoid $E(k)$. After transformations, we arrive at a standard optimal algorithm constructed on the basis of bounded ellipsoids:

$$\Theta(k) = \Theta(k-1) + \beta(k) P(k) \varphi(k) e(k); \tag{51}$$

$$P(k) = \frac{1}{\alpha(k)} \left(P(k-1) - \frac{\beta(k) P(k-1) \varphi(k) \varphi^T(k) P(k-1)}{\alpha(k) + \beta(k) \varphi^T(k) P(k-1) \varphi(k)} \right); \tag{52}$$

$$e(k) = y(k) - \Theta(k) \varphi(k). \tag{53}$$

Various modifications of the algorithm (50)–(53) can be obtained by specifying the rules for setting freely selectable parameters of the algorithm $\alpha(k)$ and $\beta(k)$. In [23], it was proposed to select the parameters as follows:

$$\alpha(k) = 1/r^2, \quad \varphi(k) = \lambda(k)/r^2. \tag{54}$$

The variable $\lambda(k)$, which is determined by minimizing the geometric dimensions of the ellipsoid $E(k)$, depends on both the step number and the values of the input variables. In principle, it is possible to accept, but this means that the newly received measures do not carry any new information. In this case, the dimensions of the ellipsoid do not change, and the estimates will not be refined.

The variable $\lambda(k)$, which is determined by minimizing the geometric dimensions of the ellipsoid $E(k)$, depends on both the step number and the values of the input variables. In principle, it is possible to accept $\lambda(k) = 0$, but this means that the newly received measurements do not carry any new information. In this case, the dimensions of the ellipsoid do not change and there will be no refinement of the estimates.

Ellipsoidal algorithms have become widespread not only in identification problems, but also in control problems when detecting disorders and monitoring parameter changes. At the same time, these algorithms, due to their numerical cumbersomeness, are poorly suited for real-time operation, and therefore their use in real-time control and management problems seems very problematic. In this regard, it seems appropriate to synthesize identification algorithms that are numerically simple, fast-acting, and suited for real-time operation in a closed loop of an adaptive control system.

3. Control systems for dynamic stochastic objects

Any control object (CO) is characterized by the following main groups of variables:

- variables characterizing the state of the object and the signals at its output, and their totality will be designated by the vector y . These variables in the control process must be maintained at a given level or changed according to a given law described by the vector y . The control accuracy may vary depending on the requirements dictated by the technology and the capabilities of the control algorithm. As a rule, the variables included in the vector are measured either directly or calculated using the mathematical model of the CO;

- variables, by changing which the regulator can affect the technological process for the purpose of control. The totality of these variables will be described by the vector of control actions;

- variables, the changes of which are not related to the impact of the control system. These variables reflect the impact of external conditions on the control object, changes in the characteristics of the process itself, etc. They are called disturbing effects and are denoted by the vector. The vector of disturbing effects, in turn, can be divided into two components, the first of which can be measured, and the second cannot. As a rule, the second component describes random effects (interference) to which any real object is subject.

A linear stochastic control object [20] can be described by the following differential equation

$$\begin{aligned}
 y(t) + \alpha_1 \frac{dy(t)}{dt} + \dots + \alpha_n \frac{d^n y(t)}{dt^n} = \\
 = \beta_0 \frac{u(t)}{dt} + \beta_1 \frac{du(t)}{dt} + \dots + \beta_m \frac{d^m u(t)}{dt^m} + w(t) + \gamma_1 \frac{dw(t)}{dt} + \dots + \gamma_l \frac{d^l w(t)}{dt^l},
 \end{aligned} \tag{55}$$

where $y(t)$, $u(t)$ and $w(t)$ are the output, control, and disturbance signals, respectively, at the time t ;
are the orders of the object by output, control, and disturbance, respectively;
are the n, m, l parameters of the object.

Qualitatively, the meaning of any problem in control theory is to select such control effects so that the control object behaves in some desired way, most often such that the output signals would be in some sense close to a given law, i.e. it is required to ensure a minimum of some function of the control error, where there is continuous time.

Control laws based on low-order controllers, such as PID controllers, can provide optimal control of objects usually not higher than the second order. Although in practice PID controllers are often used to control higher-order objects, the quality of control in this case can be far from optimal. In addition, tuning a PID controller requires high qualifications and is most often performed using heuristic rules such as the Ziegler–Nichols rule and the Takahashi rule based on it. At the same time, the importance of effective controller tuning is difficult to overestimate, since even small variations in parameters can lead to a significant change in the type of transient processes, which can lead to overexpenditure of energy resources, a decrease in the safety of technological processes and a deterioration in product quality. Thus, there is a need to develop automated procedures for tuning controllers that would reduce the influence of subjective factors, such as the experience of the operator or engineer performing the tuning and improve the quality of control. Moreover, optimal control can be achieved using controllers whose order is equal to the order of the control amplifier.

Control laws based on low-order controllers, such as PID controllers, can provide optimal control of objects usually not higher than the second order. Although in practice PID controllers are often used to control higher-order objects, the quality of control in this case can be far from optimal. In addition, tuning a PID controller requires high qualifications and is most often performed using heuristic rules such as the Ziegler–Nichols rule and the Takahashi rule based on it. At the same time, the importance of effective controller tuning is difficult to overestimate, since even small variations in parameters can lead to a significant change in the type of transient processes, which can lead to overexpenditure of energy resources, a decrease in the safety of technological processes and a deterioration in product quality. Thus, there is a need to develop automated procedures for tuning controllers that would reduce the influence of subjective factors, such as the experience of the operator or engineer performing the tuning and

improve the quality of control. Moreover, optimal control can be achieved using controllers whose order is equal to the order of the control amplifier.

Adaptive controllers can be divided into two classes.

The first class includes self-optimizing controllers, whose task is to achieve the best control quality for a given optimality criterion and the presence of certain information about the object and its signals.

The adaptation process in control systems with this type of controller occurs in three stages:

1. Identification of the object or control system as a whole.
2. Calculation of the controller.
3. Adjustment of the controller or control system as a whole.

The second group of adaptive controllers includes controllers with a reference model.

Their task is to obtain such a response of the closed control loop to a certain input signal that would be as close as possible to the response to the same signal of the reference model.

This adaptation principle assumes the presence of some measurable external signal (for example, a reference action in a tracking system), and adaptation is performed only in those periods when this signal begins to change. In this case, the adaptation process also consists of:

1. Identification of the object or the control system as a whole.
2. Calculation of the regulator.
3. Adjustment of the regulator or the control system as a whole.

The second group of adaptive regulators includes regulators with a reference model.

Their task is to obtain such a response of the closed control loop to a certain input signal that would be as close as possible to the response to the same signal of the reference model.

Conclusions

The analysis of various literary sources devoted to the solution of control problems allowed us to present the existing approaches to solving the specified problems and the place of this study in the following diagram. Here, the class of adaptive critical control systems is especially highlighted.

The problem of adaptive identification, control and critical management of dynamic, non-stationary, stochastic objects under conditions of normal operation and significant a priori and current uncertainty about disturbances acting on

the object (there is no information about the nature and character of disturbances) is considered using training models and in the presence of various types of restrictions (on amplitudes, rate of change, energy, etc.) of control and output signals, is non-stationary.

The problem of adaptive identification, control and critical management of dynamic, non-stationary, stochastic objects under conditions of normal operation and significant a priori and current uncertainty about disturbances acting on the object (there is no information about the nature and character of disturbances) using training models and in the presence of various types of restrictions (on amplitudes, rate of change, energy, etc.) of control and output signals is considered, is non-stationary. Developed in works [1–7] and summarized in the monograph [8], a new approach to the synthesis of robust methods of estimation, identification and adaptive control of static and dynamic objects under uncertainty is based on the set-theoretic approach to the interpretation of uncertainty. In this case, the main idea of the approach to the formalization of uncertainties of unknown nature is the use of parametric families of multiple estimates that cover the entire space and can be considered as fuzzy multiple estimates of unknown quantities. This idea is used for simultaneous estimation of both unknown parameters of the object and unknown characteristics of unmeasured disturbances included in the extended vector of estimates. The algorithm for adjusting the estimates of the trained model is built in the class of recurrent finitely convergent procedures for solving systems of inequalities [9].

The synthesis of control algorithms themselves is reduced to the problem of stochastic optimization of the mathematical expectation of a certain function or functional that determines the requirements for the desired value of the output signal of the object. An analysis of various literary sources devoted to solving control problems made it possible to present existing approaches to solving these problems and the place of this study in the following diagram. Here, the class of adaptive critical control systems is especially highlighted. The problem of adaptive identification, control and critical management of dynamic, non-stationary, stochastic objects under conditions of normal operation and significant a priori and current uncertainty about disturbances acting on the object (there is no information about the nature and nature of disturbances) is considered using training models and in the presence of various types of restrictions (on amplitudes, rate of change, energy, etc.) of control and output signals, is considered.

It is assumed that the parameters of the control law are a priori unknown and can change unpredictably over time, and external disturbances are of an unknown

nature (stochastic, deterministic, chaotic, etc.) and are limited in amplitude. The control goal is considered to be achieved if all inequalities are satisfied.

Thus, the work is the study and solution of the problem of adaptive critical control of dynamic non-stationary objects operating under conditions of significant a priori and current uncertainty, using training models.

To achieve the stated goal, it is necessary to solve the following problems:

- analysis and selection of structures of models, methods, procedures and spaces of signal description for solving the problem of adaptive critical control;
- development of methods and algorithms of critical control for ARMAX objects taking into account different representations of signals acting on the control object, optimization of parameters;
- synthesis of adaptive identification methods that are not based on optimization of analytical evaluation criteria and intended for operation in the circuit of a critical control system;
- search for optimal conditions that ensure maximum performance of identification methods when working with non-stationary objects;
- development and analysis of methods and algorithms of adaptive critical control based on the proposed critical control procedures and adaptive identification methods;
- simulation modeling and experimental studies of synthesized control laws for dynamic objects of various types;
- implementation of adaptive critical control methods.

References

1. Zakian V. New formulation for the method of inequalities // Proc. IEE. – 1979. – 126. – №6. – P. 579–584.
2. Avramenko V.P., Timofeev V.A. An Adaptive Robust Observation Processing System // Pattern Recognition and Image Analysis. – 1991. – Vol. 1, №1. – P. 27–28.
3. Bodyansky E.V., Vorobyov S.A., Timofeev V.A. Adaptive Recognition of the States of Dynamic Object with Periodic Output Signal // Pattern Recognition and Image Analysis. – 1999. – Vol. 9, №3. – P. 505–509.
4. Bodyansky E.V., Pliss I.P., Timofeev V.A. Discrete Adaptive Identification and Extrapolation of Two-Dimensional Fields // Pattern Recognition and Image Analysis. – 1995. – Vol. 5, №3. – P. 410–416.
5. Bodyansky E.V., Rudenko O.G. Adaptive models in control systems of technical objects. – K.: UMKVO, 1988. – 212 p.
6. Astrom K.J. Lectures on the identification problem- the least squares method Report 6806. Division on Automatic Control, Lund Institute Of Technology / – Lund; Sweden, 1968. – 164 p.

7. Ostrom K., Bolin T. Digital identification of dynamic systems based on normal operating mode data // Proc. of the II International IFAC Symposium on Self-Adjusting Systems. – M.: Nauka, 1969. – P. 99–116.
8. Soderstrom T., Stoic P. System Identification / – New York: Prentice-Hall, 1989. – 612 p.
9. Schweppe F. C. Recursive state estimation: unknown but bounded errors and system inputs // IEEE Trans. on Autom. Contr. – 1968. – 13. – №2. – P. 22–29.
10. Walter E., Piet-Lahanier H. Estimation of parameter bounds from bounded-error data: a survey // Mathematics and Computers in Simulation. – 1990. – P. 449–456.
11. Fomin V.N. Mathematical theory of trained recognition systems. – K.: Publ. KSU, 1976. – 235 p.
12. Derevitsky D.P., Fradkov A.L. Applied theory of discrete adaptive control systems. – K.: Nauka, 1982. – 216 p.
13. Yadykin I.B., Shumsky V.M., Ovsepyan F.A. Adaptive control of continuous technological processes. – K.: Naukova Dumka, 1985. – 240 p.
14. Bunich A.L. Rapidly converging algorithm for identifying a linear object with limited interference // Automation and Telemekhanics. – 1983. – No. 8. – P. 101–107.
15. Bedelbaeva A.A. Relay estimation algorithms // Automation and Telemekhanics. – 1978. – №1. – P. 87–95.
16. Parkum J.E., Poulsen N.K., Hoist J. Recursive forgetting algorithms / Int. J. Control. – 1992. – 55. – №1. – P. 109–128.
17. Kimura Y., Furukawa T., Kubota H. Generalization of the orthogonal projection algorithm into multidimensional space // Int. J. Modelling and Simulation. – 1993. – 13. – №1. – P. 35–38.
18. F.M. Al-Sadi, Petrova R.V, Timofeev V.A. Algorithms for estimating regression models in the presence of limited noise // Problems of bionics. – 2000. – Issue. 52. – P. 103–107.
19. Agadzhanov S.G., Rudenko O.G., Terenkovskii I.D. Modification of least squares algorithms for identifying parameters in the presence of limited noise // ACS and automation devices. Collection of scientific papers. – Kharkov: KhTURE. – 1997. – Issue. 105. – P. 16–21.
20. Schweppe F. C. Recursive state estimation: unknown but bounded errors and system inputs // IEEE Trans. on Autom. Contr. – 1968. – 13. – №2. – P. 22–29.
21. Al-Sunni F., Al-Nemer T. Fuzzy logic based PID self-tuning scheme for process control // Proc. IFAC-IFIP-IMACS Conf. – Belfort, France, 1997. – P. 185–195.
22. Schweppe F. S. Uncertain Dynamical Systems. – Englewood Cliffs, N.J.: Prentice Hall. – 1973. – 533 p.
23. Fogel E., Huang Y.F. On the value of information in system identification – bounded noise case // Automatica. – 1982. – 18. – No. 2. – P. 229–238.

DIGITAL TRANSFORMATION PROJECT MANAGEMENT: INNOVATIVE APPROACHES AND TOOLS

Baryshevskyi A.

In today's rapidly evolving digital ecosystem, organizations must navigate complex shifts in technology, structure, and culture. Drawing on case studies across banking, automotive, and smart city projects and reflecting on current legislative advances, this paper argues for a blended governance approach that harmonizes long-term strategy with iterative delivery. We show how AI-powered risk monitoring and embedded compliance checkpoints – illustrated by the EU AI Act's latest guidelines – enhance project resilience and stakeholder confidence. Our findings reveal a 28–33% increase in speed-to-market and marked improvements in governance transparency.

Introduction

Digital transformation transcends mere technology upgrades, encompassing fundamental changes in processes, organizational design, and corporate mindset. Unlike classic software projects, these initiatives unfold in environments where innovations such as 5G, IoT, and generative AI rapidly alter competitive landscapes and reshape user expectations. A fintech firm, for instance, pivoted its mobile payment platform twice in 2024 to integrate new biometric authentication features, aiming to outpace both regulatory demands and market entrants [1].

This volatility forces project leaders to juggle two critical imperatives. On one hand, a clear strategic vision – captured in detailed roadmaps and aligned with regulatory frameworks – anchors efforts and secures stakeholder buy-in. On the other, the capacity for swift reprioritization – enabled by regular two-week sprint cycles and real-time feedback sessions – allows teams to respond to emerging risks or opportunities without derailing overarching goals. According to recent surveys, organizations that deployed such hybrid structures reported nearly a 30% reduction in cycle time alongside enhanced sponsor satisfaction [4].

Meanwhile, the regulatory backdrop is intensifying. As of August 2, 2025, the EU AI Act enforces transparency, security-by-design, and documentation requirements for general-purpose AI models - mandates that call for compliance to be woven into project workflows from inception rather than appended at closure [2, 3]. This shift from retroactive audits to continuous governance underscores the need for frameworks that integrate compliance, risk management, and adaptive planning into every stage of digital transformation.

Features and Challenges of Digital Transformation Initiatives

Digital transformation initiatives present a confluence of technical intricacies, organizational restructuring, and cultural shift, creating a unique set of challenges for project teams. On the technical front, architects must orchestrate a symphony of public and private cloud services, legacy on-premises systems, edge devices, and AI APIs. In one multinational retailer's digital overhaul, integrating point-of-sale systems with real-time inventory management and AI-driven demand forecasting required seamless data pipelines and robust cybersecurity measures; any misalignment threatened stockouts or data breaches [15].

At the organizational level, the migration from hierarchical silos to cross-functional squads demands a redefinition of roles and responsibilities. Traditional project management offices (PMOs) often struggle to adapt, yet when a leading telecom provider shifted to a squad model for its 5G rollout, the PMO transformed into a «lean governance hub» – providing just-in-time templates, compliance checklists, and mentorship rather than rigid approval gates. This transition reduced bureaucratic lag while preserving necessary oversight [4].

Cultural transformation constitutes the third axis of complexity. In a global healthcare network's digital patient portal project, clinicians and administrative staff needed to adopt data-centric mindsets. To overcome initial resistance, project leaders implemented a «data ambassador» program, recruiting influential clinicians to champion usage, host peer-training sessions, and gather iterative feedback. This grassroots approach drove adoption from 30 percent to 75 percent active user engagement within six months [5].

Beyond these core dimensions lie secondary challenges: managing a diverse vendor ecosystem, aligning international regulatory requirements – such as GDPR, sector-specific healthcare standards, and the EU AI Act – and ensuring scalability across geographies. A cross-sector survey found that 62% of digital transformation projects encountered significant vendor integration issues, and 48% struggled with inconsistent compliance frameworks across regions, underscoring the need for a unified governance architecture capable of absorbing complexity and accelerating value delivery [15].

A Hybrid Governance Framework

Implementing a hybrid governance framework typically involves three interlinked stages, each enriched through concrete examples drawn from recent digital transformation efforts.

Strategic Blueprinting and Contextual Scoping

At the outset, executive sponsors and project managers co-create a strategic blueprint that defines long-term goals, establishes success metrics, and secures budgetary commitments. For instance, when a major Eastern European bank embarked on migrating its core systems to a hybrid cloud environment, stakeholders convened a series of multi-day workshops to align on compliance checkpoints tied to local financial regulations and the upcoming EU AI Act requirements [3]. These early engagements also yield scenario-based roadmaps, where teams anticipate potential market reactions – such as competitor feature launches – and predefine optional pivot points.

Iterative Delivery Cycles in Action

Following the strategic phase, cross-functional pods proceed with short sprints – often lasting two weeks – delivering minimum viable increments that users can test. A global automotive manufacturer, for example, used this approach when integrating an AI-powered driver-assistance system: hardware engineers, software developers, and safety regulators collaborated in each iteration to validate sensor data accuracy and interface usability. After every sprint, a demo day showcased the working prototype to regulatory liaisons and dealership managers, whose feedback directly influenced subsequent backlog refinement.

Adaptive Control Mechanisms with Live Governance Gates

To keep complex initiatives on track, organizations embed adaptive control mechanisms alongside continuous monitoring tools. In a municipal «smart city» project that deployed IoT sensors for traffic and utility management, live dashboards tracked key performance indicators – network uptime, data latency, and incident response times. When latency exceeded the agreed threshold, an automated governance gate paused further deployments until an incident review board assessed root causes and approved remediation steps. This live gating mechanism prevented service disruptions and maintained public trust.

Practical Lessons Learned

A cross-industry survey of 30 digital transformation cases reveals that projects incorporating hybrid governance:

- Accelerated critical feature releases by 20–30%, as iterative feedback reduced rework cycles.

- Lowered regulatory non-conformance incidents by 25%, thanks to embedded compliance reviews during early planning.
- Improved stakeholder engagement scores by 35%, since transparent governance gates and demos fostered shared accountability.

These practical insights underscore the value of coupling strategic foresight with executional agility – ensuring that digital transformation initiatives remain adaptable, compliant, and stakeholder-centric [11].

Advanced Risk Surveillance and Stakeholder Engagement

Leading organizations now treat risk management as an ongoing dialogue rather than a periodic checkpoint. In large-scale retail transformations – where omnichannel platforms merge online storefronts with brick-and-mortar inventory – project teams have deployed AI agents that continuously monitor sales patterns, server performance, and customer feedback streams. When an unexpected surge in mobile traffic coincided with a regional promotional campaign, the system automatically flagged a potential bottleneck in the content delivery network, prompting IT teams to scale resources in real time and prevent outages [12].

Similarly, in healthcare digitalization projects, machine learning models analyze patient admission rates, telemedicine session durations, and EHR integration errors to forecast capacity constraints. By combining these insights with stakeholder sentiment analysis – derived from clinician and patient feedback on collaboration portals – project managers can preemptively adjust staffing rosters and optimize training schedules, reducing service delays by nearly 15 percent.

Stakeholder engagement has evolved beyond routine updates. In an international energy firm’s smart grid rollout, the project office established a virtual stakeholder hub where municipal regulators, community leaders, and technical experts co-created decision criteria. Through synchronized workshops, stakeholders reviewed live risk dashboards, debated mitigation options, and signed off on deployment phases collectively. This participatory model not only accelerated approvals but also fostered trust, which proved crucial when unexpected regulatory changes required midstream design adaptations.

These cases illustrate how embedding AI-driven surveillance tools and collaborative engagement frameworks transform risk from a roadblock into an opportunity for alignment and innovation. By continuously integrating data-driven foresight with multi-stakeholder input, digital transformation initiatives navigate complexity while maintaining organizational momentum.

AI-Powered Decision Support and Compliance Automation

The infusion of AI into project management extends beyond risk surveillance. Natural language understanding engines can autonomously generate progress narratives, tag documentation by topic, and surface regulatory clauses requiring attention. Decision support modules, drawing on optimization algorithms, recommend resource redistribution to address forecasted bottlenecks in real time.

To safeguard ethical and legal integrity, continuous compliance automation integrates AI Act obligations directly into DevOps pipelines. At each build and deployment stage, automated scripts verify model training data provenance, assess bias metrics, and ensure security-by-design controls. Audit logs capture version histories, model performance statistics, and incident responses – creating an immutable trail for regulators and auditors [3]. Firms adopting such measures report a 22 percent reduction in post-deployment non-conformities and streamline certification processes for new AI capabilities [9].

Conclusion and Future Directions

Navigating the ever-shifting landscape of digital transformation demands that project governance models evolve beyond static frameworks. Our analysis underscores the imperative of embedding adaptability and compliance into the core of project lifecycles. By weaving iterative feedback loops into strategic planning, organizations gain the ability to course-correct in real time, mitigating emergent risks before they escalate. The empirical cases from banking, automotive, and smart city initiatives illustrate how hybrid governance not only accelerates feature delivery but also fortifies stakeholder trust through transparent decision gates [11].

Looking ahead, the integration of AI into governance structures presents both opportunities and complexities. Generative AI tools have the potential to automate aspects of roadmap refinement and stakeholder communication, yet they introduce new risk vectors – such as model hallucinations and data privacy concerns – that must be preemptively managed. Future research should therefore investigate how AI-driven governance assistants can be trained under compliance constraints and audited for ethical performance [13].

Moreover, scaling these frameworks across small and medium-sized enterprises (SMEs) remains an open challenge. SMEs often lack the resources to deploy sophisticated AI surveillance platforms or maintain dedicated compliance teams. Pilot programs could explore lightweight hybrid models that leverage

cloud-based governance-as-a-service offerings, enabling SMEs to adopt best practices without prohibitive upfront investments [14].

Finally, longitudinal studies are needed to assess the long-term impact of regulatory regimes like the EU AI Act on innovation velocity. While immediate compliance measures might introduce overhead, iterative adjustments to act frameworks could yield a balanced ecosystem that supports both risk mitigation and accelerated technology adoption. Such research will inform policymakers and practitioners striving for governance architectures that evolve in step with technological progress.

References

1. Petrenko, V.O., & Baryshevskyi, A.I. (2024). Features of digital transformation project management. *Journal of Digital Project Studies*, 5(2), 34–49.
2. European Commission. (2025, July 30). The second enforcement deadline for the EU AI Act is approaching. *IT Pro*. Retrieved from <https://www.itpro.com/business/policy-and-legislation/the-second-enforcement-deadline-for-the-eu-ai-act-is-approaching-heres-what-businesses-need-to-know-about-the-general-purpose-ai-code-of-practice>
3. European Parliament. (2024). Regulation (EU) 2024/1689 on artificial intelligence (EU AI Act). *Official Journal of the European Union*.
4. Baryshevskyi, A.I., & Petrenko, V.O. (2025). Methods of digital transformation project management under conditions of instability and rapid technological change. *European Journal of Project Management*, 12(1), 12–29.
5. Protiviti. (2025, February). *Whitepaper: EU AI Act – Compliance and Best Practices*. Retrieved from Protiviti website.
6. Financial Times. (2025, July 29). Google set to sign EU code of practice on development of AI. Retrieved from <https://www.ft.com/content/97ee867c-64c2-44ee-a519-36ca670ed565>
7. CIO.com. (2025, July 31). EU AI Act: one year on, new measures enter effect. Retrieved from <https://www.cio.com/article/4032894/analysis-of-the-european-ai-regulation-one-year-after-its-entry-into-force.html>
8. Gartner. (2025). *Top Trends in Digital Project Management*. Retrieved from Gartner research portal.
9. Sembly AI. (2025). *EU AI Act Overview and Impact on AI Tools*. Retrieved from <https://www.sembly.ai/blog/eu-ai-act-overview-and-impact-on-ai-tools/>
10. BSK Legal. (2025, June 12). The EU AI Act: What U.S. Companies Need to Know. Retrieved from <https://www.bsk.com/news-events-videos/the-eu-ai-act-what-u-s-companies-need-to-know>
11. McKinsey & Company. (2025). *Accelerating Digital Transformation in Financial Services, Automotive, and Smart Cities*. Retrieved from <https://www.mckinsey.com/industries/financial-services/our-insights/accelerating-digital-transformation>

12. McKinsey & Company. (2025). Real-Time AI Surveillance in Retail and Healthcare Projects. Retrieved from <https://www.mckinsey.com/industries/retail/our-insights/real-time-ai-surveillance>
13. Smith, J., & Lee, A. (2025). Training AI Governance Assistants: Ethical Frameworks and Compliance. *International Journal of AI Ethics*, 3(1), 45–62.
14. Lopez, M., & Zhang, T. (2025). Governance-as-a-Service for SMEs: Enabling Digital Transformation at Scale. *SME Tech Journal*, 8(2), 112–131.
15. Gartner. (2025). Overcoming Vendor Integration and Multi-Region Compliance in Digital Transformation Projects. Retrieved from Gartner research portal.

INTEGRATION CAPABILITIES FOR ENTERPRISE VALUE CHAIN MANAGEMENT

Bulavin D., Petrenko V.

The widespread use of decision support systems has become an indispensable tool for managing complex value chains to the customer. We explore integration processes that encompass both technological and organizational integration. Renowned methodologies in enterprise integration state that integration processes need to be planned. However, planning cannot be precise as business processes and technologies are changing rapidly, offering different potential solutions. Integration approaches can be used to create flexible, loosely coupled structures that can quickly meet specific integration needs. These new technical means depend on standards that work both inside and outside the enterprise.

Introduction

The current state of economic development of Ukraine requires the study of effective ways of information support of goods at all stages of the life cycle. An important aspect of the activity of enterprises is the creation of an integral information environment. In the professional literature, the integration of the enterprise is usually considered at two levels – the system level and the organizational level. Integration at the system level requires common standards and definitions of data, as well as certain means of synchronization of communication between different software applications. Usually, this is what is meant by the recently coined term «Enterprise Application Integration». However, as Marcus and others have pointed out [1], system integration in the sense of a software system alone is not sufficient to ensure organizational efficiency and effectiveness. Organizations are made up of individuals, departments, divisions, and functions that must also be integrated for the success of the organization. Today, integration is a more generalized term that refers to both behavioural patterns focused on both people and systems. However, integration is most often used when discussing relationships between software systems. According to many researchers, the integration of project management methodology into management processes provides companies with a real chance to provide both organizational and technological support for the implementation of the enterprise development strategy.

The purpose of this study is to understand the experience of practitioners of both disciplinary approaches and identify common mechanisms for their implementation.

Review of scientific literature on the topic

The publications were analyzed according to both strands of modern scientific research, which led to a number of methodological recommendations for potential scientists. Two main provisions can be observed here: firstly, digital technologies provide customers with the opportunity to create products together with the manufacturer, for example using digital platforms [2], and secondly, at the external level, value is increasingly being created from a series of partnerships in a network of shared values [3] and others [4], define coordination as the management of dependencies between all activities of the organization. Many academic researchers propose to apply a systematic technical approach to the integration of enterprises. From the point of view of software developers, integration has little to do with the specifics of the business. Instead, it is a question of how to understand the value creation processes of a particular business. First, Mary Shaw [5], noted that the issue of interconnection of systems was often not sufficiently thought out in relation to integration processes. But they should be the main issue in the technological renewal of the enterprise. Cris Dellarocas of the Massachusetts Institute of Technology pointed out [6] at one time that there may be a need to give first-class status to component interdependencies, since interdependence is a higher-level concept than a simple relationship.

Unfortunately, current production management systems cannot support comprehensive solutions for all scenarios. Most business applications are legacy systems that have been developed over the years. While each of these software systems does a good job with its assigned tasks, they are not equipped to handle business scenarios across multiple application domains. The growing need for customized products and services across many industries has made today's global supply chains more complex. This increased level of supply chain complexity, in turn, has increased the level of uncertainty and risks faced by modern companies [7]. To reduce this level of uncertainty, businesses are trying to manage flexible supply chains in a way that allows them to respond to demand [8, 9]. A large number of different definitions of enterprise supply chain flexibility have been developed using conceptual models, normative indices, and structural modeling [10]. Traditional systems do not provide the right solutions for flexible supply chains. Their design ideology does not provide the high level of decentralized governance required for flexibility [11]. They are characterized by inflexibility in terms of supply chain reconfiguration, high development and maintenance costs, and limited computing resources to manage complex systems [12]. According to many researchers, the integration of information technologies into management processes provides

companies with a real chance to provide support for the implementation of the strategy of digital transformation of the enterprise.

Presentation of the main material

In the space of growing chaos, it is already clear to everyone that as corporations grow, the need for integration in the enterprise also grows. Integration is really a business need, and it affects not only technological mechanisms, but also requires reorganization (renewal) of proper organizational structures, goals and incentives. Table 1 shows the aspects of technical and organizational integration with a description of the known integration mechanisms. The list of resources and activities required is shown in the first column, examples of software and means of coordinating people in the second column, and in the third column – elements of supporting infrastructure – in the right column.

Table 1

Analysis of technical and organizational integration of the enterprise with a description of known tools

Resource Integration Need	Integration Tools	Infrastructure environment
1. Organizational Units (Functions/Departments)	E-mail, collaborative software, lateral teams, strategy, budgets, performance metrics	Organization policies/ structure
2. Decision Makers	Email, collaborative software, knowledge management systems, face-to-face meetings, job design, performance metrics	
3. Business Processes (both internal & external)	Workflow, Collaborative Systems, SCM, CRM, Web Services, process owners, teams, performance metrics, service level agreements	Systems Architecture
4. Applications	Inter-process communication, Messaging, ERP, Web Services, Organizational Integration	
5. Data	Data Dictionaries, Databases, XML	

From the analysis of the table, it is concluded that effective integration requires attention to the elements at all levels in both horizontal and vertical dimensions. An integrated architecture exists to maintain coordination in the use of the firm's material, financial and human resources. Integration levels are described from top to bottom – from the most general to the most specific. It should be noted that integration is necessary not only within different layers, but also between layers.

Our research focuses on the central column of the table, which also mentions the role of infrastructure elements – standards, architecture, networks and organizational structure – in promoting integration both within and between layers.

The upper level of organizational integration of enterprises is focused on the layers of software tools within the existing structure along with the architecture and standards of the network infrastructure. Professional literature on organizational integration of enterprises is mainly focused on the levels of decision-making and integration of various departments with an auxiliary organizational structure and development strategy. To achieve this strategy, the integrated technical and organizational components, must coordinate the financial and human resources of the organization to achieve the objectives of the organization.

As you know, any organization is divided into several departments or departments [13]. This division occurs because each department of the organization must focus on the fulfillment of certain functional responsibilities. This specialization leads to different attitudes of managers to the triangle of achieving goals: orientation to time, costs and quality. Organizational integration requires, firstly, the establishment of communication in such a way that departments are informed about the activities of other departments and cross-functional processes are carried out smoothly for the sake of common higher-level goals.

Proponents of the reengineering movement proposed matrix organizational structures in which functional tasks are subordinated to a process approach. Adding process owners and cross-functional teams to the organization promotes horizontal coordination, but requires some kind of matrix management. The main challenge in achieving departmental coordination is that the vertical chain of command and authority must be aligned with the needs of the organization's horizontal processes. To address this problem, Rammler and Brahe [14] recommend that the share of resources and goals of each department be determined based on the department's contribution to cross-functional processes. This approach differs significantly from the usual practice, in which functional goals and performance indicators are set by functions together with senior management without special consideration of the needs of horizontal coordination.

The second layer of integration concerns decision-makers. The integration of these individuals usually occurs when they coordinate their actions for the benefit of the enterprise and are free to share their knowledge. To achieve this integration, strategic decisions must be made by the highest levels of management. While classic models of organizational hierarchies assume that top-level leaders are isolated at the top, more recent publications show that decision-makers are in constant contact

with each other, and their interactions are critical to the success of the company. A newer way of thinking has led to the recognition that interactions between executives in a company are not just about tasks. The psychological side of integration in organizations should not be underestimated, since there are many examples in which large projects have failed due to the interpersonal frictions of higher-level managers. It is for this reason that decision-makers opt for less formal forms of communication with their colleagues. And it is for this reason that companies encourage social interaction between their executives.

Computer systems, as a rule, provide decision-makers with concise rather than extensive information. Multimedia tools definitely increase the availability of information. To promote communication, special software tools have been developed to help decision-makers coordinate their activities, communicate and exchange information and ideas.

Integration at the business process level consists of related sets of actions that are performed by people and software agents according to business rules that can be applied more or less strictly. The communication between the software and the people who carry out the process is well integrated or coordinated when the process is efficient, precise and relevant to the task at hand from a technical and organizational point of view. In addition, the objectives of each workflow must be aligned with those of the organization as a whole. This is called intra-process integration. External process integration means that an organization can seamlessly link its internal processes with those of its suppliers, intermediaries, and customers. A distinction is made between automated and partially automated process integration. Fully automated internal/external process integration is considered achieved when a firm connects to external agencies in real-time. At a higher level of integration, external processes must be coordinated between all firms in the value chain to achieve improved productivity and service.

Workflow management systems are specifically designed to support the automation of business processes by moving work between participants according to certain rules [15]. These systems provide visual interfaces for process design, manage process instances, and easily interact with an organization's legacy, client-server, and ERP applications. They also contain simple organizational models, such as who reports to whom, who has the authority to approve what, who can perform which roles, and who has access rights to which data and which applications. And they perform the function of coordinating resources, managing work assignments and load balancing between process participants. In this sense, the workflow promotes integration at all levels, as indicated in Table 1. Historically, workflow has been used

primarily to coordinate cross-functional and internal processes. Right now, the main trend in workflow improvement is focused on integrating the entire value chain.

The fourth layer of integration considers the integration of software applications. When application sets are integrated, each of them can cause the other to perform a function. The goals of software integration are:

1. Break down complex processing tasks into separate steps performed by relatively independent programs.
2. Support the interactive operation of universal shopping services by combining different applications in a common interface.
3. Ensure the integration of data and information.

Application integration is now supported by a wide range of software. Initially, such integration was supported by calls to remote procedure call protocols, then transaction managers, and now application servers. A higher level of integration is demonstrated by ERP systems connected through a common database. But ERP systems cannot always bridge the gap between application and process layers. This gives rise to complaints that it is easier to adapt an organization to an ERP system, than to adapt the ERP system to meet organizational requirements [16].

At the lowest level of integration, data integration is defined. The goal of data integration is to allow organizations to combine, aggregate, and report on data from different sources. At the syntactic level, the software can process data stored on different devices and in different formats. At the semantic level, the values of all data items should have the same meaning across applications. Currently, the most popular solution for data integration is the use of Extensible Markup Language (XML). XML combines data and its description in one place, allowing applications to access and understand data stored in heterogeneous files and different databases.

It was established that the integration of the enterprise must take place at many technical and organizational levels. Some techniques that can be used at each level have been introduced. Every organization needs an information system architecture that supports enterprise-wide integration. If the enterprise is fully integrated, then data from all existing applications can be shared. Such an ideal of integration has existed throughout the history of the development of information systems. But the search for the perfect integrated enterprise system continues. General frameworks are preferable to point solutions. Now enterprises are striving to solve the problem of integration as a whole, and not with separate point solutions. This encourages developers to design systems that cover all business processes of enterprises. Currently, integration mechanisms tend to weaken connections

between the elements of the system, since in a loosely coupled environment, it is easier to make a change in one separate part of the system.

The use of the XML data description language is becoming the de facto standard in the enterprise. This is important because XML can describe all data transferred between systems in the same way; significantly reducing the many file formats supported by enterprises. Industry-wide efforts stimulate internal efforts. So, for example, a set of definitions of data used by financial services companies to interact with each other can also be used internally to integrate corporate systems. After that, internal and external communication can use the same standard set of messages. Both standards focused on specific manufacturing industries and the new ISO standard focused on the financial industry are examples of industry models that are reflected in enterprises.

Consider three approaches to integration that have been proposed in both industry and academia. The first approach is sequential integration of $A \rightarrow B$. Over time, this kind of sequential integration between applications has become more complex with the advent of different programming languages. More modern approaches to sequential integration, such as the UNIX shell system, provide a user-friendly interface, but also have certain limitations. That two different applications can exchange information using the same database server. This approach can solve problems that consistent integration cannot handle. For example, any transactional system, such as a booking system, can be implemented using database technology. The third approach considers each part of the system as a separate entity that can either create or receive messages [17]. This approach can work at any scale files exchanging between large information systems.

Taking into account the above approaches, it is possible to describe commonly used technical mechanisms of integration on an enterprise scale. Most companies do not integrate two technological systems until they absolutely do not need it. The easiest way to integrate is to set up a point-to-point interface. The order management system can generate a file every night that is read by the commission system the next morning. The disadvantage of this mechanism becomes obvious with the growth of the number of systems in the enterprise. An IT organization can be overwhelmed by the need to update and maintain file interfaces. The incremental cost of each new interface is low, but the cumulative running costs for the organization are high.

Many companies can integrate a significant portion of its operations by replacing financial, HR, and manufacturing systems with an ERP system provided by one of several major software vendors. All applications will use the same data model,

the same databases, and the same interfaces. This approach has worked for many companies. But ERP is more of an approach than a technology, and there are indications that the next generation of ERP systems will take advantage of new technological mechanisms, such as the Value Chain Management System (VCMS). VCMS is a framework that helps organizations manage all activities related to the creation and delivery of a product or service, from the initial phase to the final delivery to the customer. It focuses on optimizing every step of the process to add maximum value to both the business and the customer.

Table 2 summarizes the capabilities and advantages of modern IT tools for effective VCMS. It should be emphasized that the autonomous decision-making powers of information systems, combined with built-in capabilities for big data analytics, can significantly increase the flexibility of the value chain.

Table 2

**Capabilities and Advantages of Modern IT Tools
of Value Chain Management Systems**

Element	Ability	Overview
1	2	3
Information technology and integration	Capture demand information immediately	An agent can interact with other agents or humans through the use of an agent communication language
	Information accessible to supply chain wide	The realization of a holistic cross – organizational collaboration is possible. The short and simple nature of software is an important enable of this future.
	Integrate with Semantic Web Service	Overall, the use of wrapper agents can significantly facilitate integration with semantic web services. The leverage of semantic web services leads to a higher degree of information integration.
Customer / marketing sensitivity	Perceive opportunities to increase customer value	Can diagnose new opportunities that could enrich customer value. Genetic algorithms leverage seems an important facilitator.
	Customer-driven products	Can robustly response to customer preferences due to their inherent learning capabilities
	Incorporate Semantic Web Services	Overall agents through their learning ability can leverage perceptions taken from semantic web services in order to achieve a higher level of customer sensitivity

Continuation of the Table 2

1	2	3
Process Integration & Performance management	Facilitate Rapid decision making	Can demonstrate real – time responsiveness capability. In fact, learning ability is an important facilitator of this capability. Moreover, they can be reconfigured with easiness to new business processes
	Pro-actively update the mix of available supply chain process	Can real-time perceive their environment (software agents, manufacturing equipment). In this manner, they can real time adjust manufacturing processes.
	Leverage Semantic Web Services	Overall, performance management can be enriched with big data interpreted under a meaningful way.
Collaborative planning	Representation of trust-based relationships with suppliers/customers	Based on previous research trust among partners can be represented among agents matching reality.
	Reduce Bullwhip effect – data accuracy	Agents through the use of learning algorithms can achieve solutions that approach to optimal sharing
	Leverage Semantic Web Services	The data inference can provide meaningful information for all the nodes of the chain.

Conclusions

The analysis makes it possible to visually assess integration management strategies in organizations that carry out project activities. We discussed which technologies are suitable for software applications in organizations where shared value chain processes are required. Thus, we theoretically linked big data analytics and value chain performance. However, we understand that our study lacks empirical research in the application of the proposed frameworks. ensures that our provisions are limited. However, our further research will be aimed at identifying and validating the potential benefits of Value Chain Management Systems. Future research may also focus on building a general theory of organizational integration. For example, the question of how certain qualities of value chain managers affect the implementation of human-centered technologies in the decision-making process would be useful in identifying the propensity of potential organizations to use Value Chain Management Systems as a tool to support decision-making.

References

1. Markus, M.L., Paradigm Shifts – E-Business and Business/Systems Integration. Communications of AIS, 2000. 4(10).
2. El Sawy OA, Kræmmergaard P, Amsinck H, Vinther AL (2016) How LEGO built the foundations and enterprise capabilities for digital leadership. MIS Q Exec 15(2):141–166.

3. Evens T (2010) Value networks and changing business models for the digital television industry. *J Media Bus Stud* 7(4):41–58
4. Malone, T.W., et al., Tools for inventing organizations: Toward a handbook of organizational processes. *Management Science*, 1999. 45(3): p. 425–443.
5. Shaw, M., Procedure Calls are the Assembly Language of Software Interconnection: Connectors Deserve First-Class Status, in *Studies of Software Design, Proceedings of Workshop, Lecture Notes in Computer Science No. 1078*, D.A. Lamb, Editor. 1996, Springer-Verlag.
6. Dellarocas, C., A Coordination Perspective on Software Architecture: Towards a Design Handbook for Integrating Software Components. 1996, MIT.
7. Manuj, I., and Mentzer, J., 2008. Global supply chain risk management strategies. *International Journal of Physical Distribution and Logistics Management* 38 (3), 192–223.
8. Braunscheidel, M.J., and Suresh, N.C. 2009. «The Organisational Antecedents of a Firm's Supply Chain Agility for Risk Mitigation and Response». *Journal of Operations Management* 27 (2), 119–40.
9. Agarwal, A., Shankar, R. & Tiwari, M. (2007). Modeling agility of supply chain, *Industrial Marketing Management*, 36:443-457.
10. Gligor, D. M., & Holcomb, M. C. (2012) «Understanding the role of logistics capabilities in achieving supply chain agility: a systematic literature review» *Supply Chain Management: An International Journal*, 17 (4), 438–453.
11. Mayer-Schönberger, V., and Kenneth, C. (2013) *Big data: A revolution that will transform how we live, work, and think*. Houghton Mifflin Harcourt.
12. Akkermans, H., Bogerd, P., Yücesan, E. & van Wassenhove, L.N. (2003) «The impact of ERP on supply chain management: exploratory findings from a european delphi study». *European Journal of Operational Research* 146 (2), 284–301.
13. Lawrence, P.R. and J.W. Lorsch, *Organization and environment; managing differentiation and integration*. 2007, Boston,: Division of Research Graduate School of Business Administration Harvard University.
14. Rummler, G.A. and A.P. Brache, *Improving performance : how to manage the white space on the organization chart*. Jossey-Bass management series. 1990, San Francisco: Jossey-Bass.
15. Wfmc, *The Workflow Reference Model, Version 0.6.*, <http://www.wfmc.org/>
16. Esteves, J. and J. Pastor, *Enterprise Resource Planning Systems Research: An Annotated Bibliography*. Communications of the Association for Information Systems
17. Rumbaugh, J., I. Jacobson, and G. Booch, *The unified modeling language reference manual*. The Addison-Wesley object technology series. 1999, Reading, Mass.: Addison-Wesley.

EXPERIMENTAL EVALUATION OF MODERN UI/UX DESIGN TOOLS AS A COMPONENT OF PROJECT MANAGEMENT

Cherkasov D., Kuznetsova Yu.

In today's digital transformation environment, the quality of UI/UX design of web applications determines the competitiveness of the product and the level of user satisfaction. The research problem lies in the insufficient formalization of methods for assessing the effectiveness of prototyping tools and interface implementation, which complicates the selection of optimal technologies in the process of IT project management. The article presents an experimental model aimed at comprehensively evaluating modern UI/UX design tools, in particular prototyping environments (Figma, Adobe XD) and CSS frameworks (Bootstrap, Tailwind CSS, Material UI). The study is based on quantitative (task execution time, number of errors, code volume, productivity) and qualitative (SUS, NPS scores, click heat maps, UX reviews) analysis methods. The experimental results showed that Figma provides 17% faster prototyping compared to Adobe XD, while Tailwind CSS demonstrated the highest interface performance according to Lighthouse indicators. Material UI, in turn, confirmed the effectiveness of using Material Design principles to improve usability. The results obtained can be integrated into the project management process to make informed decisions regarding the choice of UI/UX design tools. Thus, the research contributes to the development of innovative IT project management technologies and improves the quality of digital products.

Introduction

Relevance of research. In today's digital environment, UI-UX design is one of the main aspects in creating web applications, which must not only perform their functions, but also be intuitive and user-friendly. Users expect high-quality interaction with interfaces that provide fast and efficient work, which affects the overall success of the product in the market.

Given the large number of competitive solutions, web applications that offer a poor user experience have a significantly lower chance of success. A well-thought-out UI-UX design can not only increase user engagement, but also reduce bounce rates and increase conversion [1]. The constant development of technologies and tools for creating interfaces requires the exploration of new approaches, methods and tools for UI-UX development to adapt to changing user needs.

An analysis of publications by foreign scholars has shown that the subject area of UI-UX design in the context of web development remains insufficiently formalized. There are a large number of methods and approaches that offer general recommendations, but they often do not reveal key aspects specific to web applications. For example, publications by researchers such as Alan Cooper

(author of the book «The Inmates Are Running the Asylum» [5]) and Norman Donald (author of «The Design of Everyday Things» [7]) offer general principles of UX design, but specific recommendations for web applications are lacking.

Among the methods that are currently actively used in UI-UX design of web applications, the following can be distinguished:

- Design Thinking [3] is a popular approach used by companies such as IBM and IDEO to develop innovative products, including web applications. However, this method focuses more on the general aspects of human interaction with a product and does not offer specific tools for web applications.

- Human-Centered Design (HCD) [4] is a methodology that puts the user at the center of the development process. Although this method is mainstream at many companies such as Apple and Microsoft, its specific guidelines often need to be adapted to the modern needs of web applications.

- Atomic design [2] is a methodology in UI design that allows for the optimal design of a design system. The concept was formulated by interface developer Brad Frost under the slogan «Build systems, not pages». In other words, instead of designing, for example, an interface page in its entirety at once, it is more logical to assemble it in parts.

UI-UX issues in web design are also addressed by scholars such as Steve Krug (author of the book «Don't Make Me Think») [6], who emphasizes the simplicity of interfaces but provides few specific recommendations regarding modern web development technologies.

Companies like Figma, Adobe, and Sketch offer advanced tools for creating UI-UX designs, but most of their solutions address general interface design issues and do not address deeper problems specific to web applications. This situation highlights the need for detailed experimental research into modern methods, tools, and technologies that can be adapted or optimized for creating effective UI-UX design for web applications. Studying such approaches will allow us to identify the most relevant and effective technologies for ensuring user interface usability in the context of modern web development.

Thus, this study aims to address this gap by offering specific recommendations and tools to improve the web interface development process, which will significantly improve the quality of user interaction with web applications and contribute to the development of digital products in various industries.

1. Setting the goals and objectives of the experimental study

General characteristics of the research methodology. The modern development of digital technologies places increased demands on the quality of web application interfaces, where not only visual appeal, but also the effectiveness of user interaction plays a key role. This section presents the methodology of an experimental study aimed at a comprehensive assessment of different approaches to creating UI/UX design. The study will cover a comparative analysis of modern prototyping tools (Figma, Adobe XD), interface implementation frameworks (React, Tailwind CSS, Bootstrap, Material UI), and design methodologies (Design Thinking, User-Centered Design, Atomic Design).

The main focus of the research will be to identify the most effective approaches to interface design through a comprehensive evaluation of various combinations of tools, methodologies, and qualitative user assessments (SUS, NPS), which will allow obtaining substantiated conclusions regarding the effectiveness of individual approaches in specific usage scenarios.

The purpose of this experiment is to determine the most effective methods, tools and technologies for developing UI/UX design for web applications in terms of usability, user productivity and speed of implementation.

To achieve the goal, you must complete the following tasks:

- analyze feedback and experiment factors;
- compare popular UI frameworks and libraries;
- determine the impact of different design tools on interface quality;
- empirical testing of the effectiveness of created interfaces using UX metrics;
- describe the task of the experiment;
- choose a method for processing the experiment results;
- determine the technical means of implementation;
- determine the software implementation tools.

2. Procedure for experimental studies methods, tools and technologies

Experimental research is implemented based on a structured approach that includes three main components: methods for evaluating user experience, tools for creating interface prototypes, and technologies for implementing web interfaces.

At the first stage of the study, the objects of the experiment are determined, which include:

- UX evaluation methods: System Usability Scale (SUS), Customer Effort Score (CES), Net Promoter Score (NPS), A/B testing;

- tools for creating interfaces: Figma and Adobe XD;
- interface implementation technologies: Bootstrap, Tailwind CSS frameworks and Material UI UI library.

For each object, the corresponding measured parameters are determined, which allow for an objective comparison.

The experimental study of the methods is described below:

- System Usability Scale (SUS). Users are asked to fill out a standard questionnaire of 10 questions, where they rate the interface on a scale from 1 to 5. Based on the answers, an average SUS score is calculated, which shows the level of usability of the interface. Input data: SUS questionnaire; Output data: average score;
- Customer Effort Score (CES): Users rate how easy it was to complete a task (such as signing up) on a scale of 1 to 7.
- Net Promoter Score (NPS). Users answer the question: «How likely are you to recommend this app to others?» on a scale of 0 to 10. The results are categorized into promoters (9–10 points), passives (7–8 points), and detractors (0–6 points). Input: NPS question; Output: net NPS score;
- A/B testing. Users are presented with two interface options, such as different element designs or placement. The attractiveness score is measured on a scale of 1–10, and data is collected from open comments, as well as the percentage of users choosing each interface. Input: two interface options; Output: percentage of users who preferred option A or B, average score on a scale of 1–10 for each option.

These indicators are recorded through the Lyssna platform and Google Forms, and assessments are collected both within the Prototype Test and Preference Test, and through questionnaires after completing the scenario, to collect quantitative data on user interaction with the web application, in particular, these include: task completion time, number of errors, task completion rate.

Experimental research of Figma and Adobe XD tools is evaluated in terms of the speed of creating prototypes, the number of clicks/actions to the final version, and user interviews are additionally recorded regarding the ease of use of each tool, which allows obtaining high-quality UX feedback:

- input data: the same tasks for both tools;
- Initial data: execution time, number of iterations, participant feedback.

Experimental research of Bootstrap, Tailwind CSS and Material UI technologies, evaluated by the amount of CSS code (in KB), the speed of interface implementation (in hours) and overall performance (Lighthouse scores). In parallel,

a qualitative assessment is carried out: compliance with guidelines, ease of customization of elements:

- input data: interface layout;
- source data: comparative data for each criterion (CSS code volume, creation speed, performance).

All recorded parameters are divided into quantitative and qualitative assessment methods.

Quantitative methods include:

- task completion time;
- number of errors;
- convenience score (SUS);
- number of clicks to the target action;
- CSS file size in KB;
- interface implementation speed (in hours).

Qualitative methods include:

- UX user feedback;
- interaction heatmaps;
- compliance with guidelines;
- post-test interview.

At the final stage, the data undergoes two-stage processing:

- primary processing involves aggregating data from Lyssna and Google Forms, calculating average task completion time, number of errors, scores on SUS, CES, NPS scales, A/B testing results (via Preference Test). In addition, primary processing includes quantitative indicators obtained during experimental research of interface implementation using various technologies (Tailwind CSS, Bootstrap, Material UI), in particular, the volume of generated CSS code (in KB), interface development speed (in hours) and implementation performance according to Lighthouse indicators;

- Secondary processing includes correlation analysis, visualization of results (diagrams, graphs, histograms), as well as generalization of conclusions about the effectiveness of each method, tool, and technology.

The input data at this stage are respondent ratings, clicks, time, satisfaction level, heat maps, etc. The output results are structured performance tables, correlation graphs, and a comparative conclusion on the optimal combination of tools and technologies for creating UX-oriented interfaces. If there are not enough tools for conducting experiments, a custom script will be created.

Additionally, the reliability of the obtained results is ensured by repeating the same test scenarios on several user groups, which allows reducing the influence of subjective evaluations. To increase the accuracy of conclusions, statistical analysis methods such as standard deviation and confidence intervals are also applied. Comparative experiments are conducted in controlled conditions, which provides reproducibility of the study. Moreover, qualitative data from user interviews are triangulated with quantitative indicators, ensuring consistency of research outcomes. Finally, the combination of several UX evaluation methods in parallel (SUS, CES, NPS, A/B testing) increases the robustness of the conclusions and allows the detection of potential contradictions between different types of metrics.

Fig. 1 presents a diagram of experimental research methods, tools and technologies, for a better understanding of the sequence of actions and the relationship between the stages of the study.

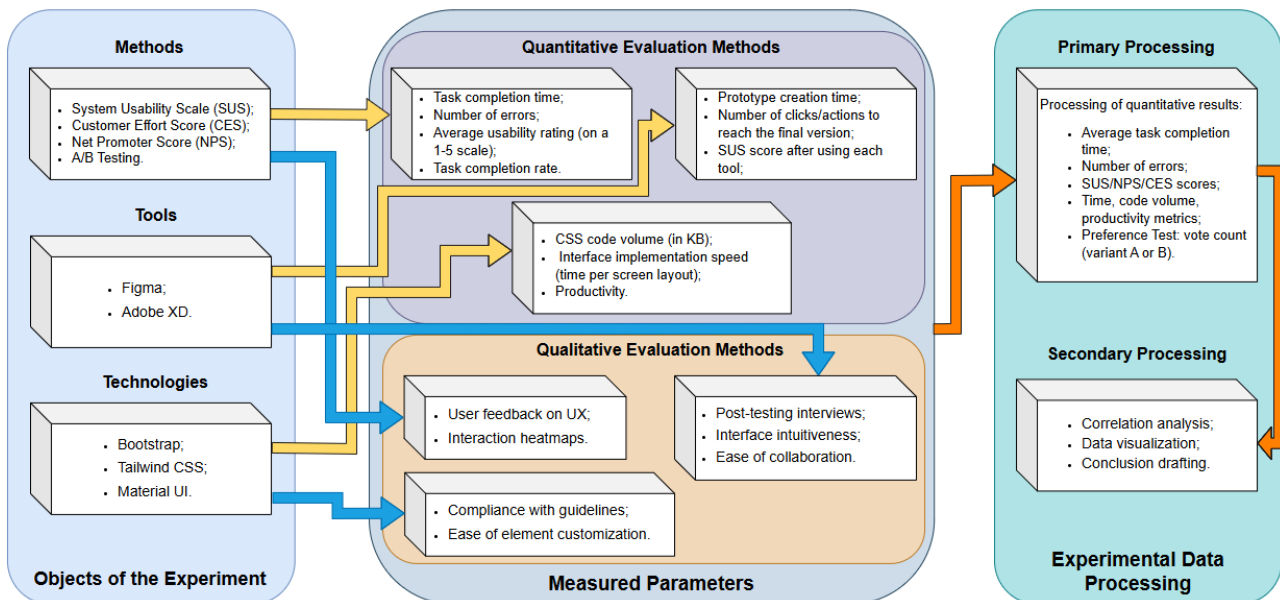


Fig. 1. Scheme of experimental research

3. Analysis of factors and feedback of the experiment

Basic principles of the theory of experimental design. Experimental design in UI/UX research is based on general principles of the scientific approach, adapted to the specifics of digital interfaces. The theoretical basis of such planning involves the consistent identification of key components of the research process that ensure its objectivity and reproducibility.

The first stage is the analysis of existing research in the field of interface design. This step allows you to identify relevant research directions, avoid

duplication of already obtained results and form an effective methodology for conducting your own experiment. Of particular importance is the study of similar works that relate to the comparison of design tools and development frameworks.

Selecting dependent variables, or feedback, for an experiment is a critical planning step. In the context of UI/UX research, feedback should adequately reflect the effectiveness of the interface and the quality of the user experience. Each selected indicator should be quantifiable, unambiguously interpreted, and have clear practical meaning. For example, task completion time or the number of user errors are objective metrics, while subjective usability assessments require additional attention to the methods of their collection and analysis.

Independent variables, or factors of the experiment, define the main parameters to be studied. In our case, these include the design tools used, development frameworks, and interface creation methodologies. Unlike feedback, the number of factors can be relatively large, but their selection should be determined by a clear understanding of the research objectives and practical use [8].

An important aspect of planning is determining the parameters that will be held constant throughout the experiment. These decisions are made based on the research objectives and available resources, allowing you to focus on the most relevant aspects of interface design.

Determination of experimental factors. When planning an experimental study in the field of UI/UX design, special attention should be paid to identifying key factors that will influence the quality and effectiveness of interfaces. The selection of these factors should be based on a clear understanding of their potential impact on the final result. An insufficiently thorough approach to identifying factors can lead to incomplete or biased results, as the failure to take into account important variables can significantly distort the picture of the study.

The factors of this study will be:

- prototyping tools (Figma and Adobe XD);
- framework and implementation library (Bootstrap; Tailwind CSS, Material UI);
- design methodology (Design Thinking, UCD, Atomic Design).

Determining the feedback of an experiment. When planning an experiment, a key step is to determine the parameters that will serve as criteria for evaluating the effectiveness of the solutions under study. These parameters, which are called responses or optimization parameters in scientific terminology, should clearly reflect the objectives of the study and be quantitatively measurable.

Real objects and processes are characterized by high complexity, which necessitates the consideration of a set of interrelated parameters. Each object can be

characterized by the entire set of parameters, or by one – a single optimization parameter. The remaining characteristics do not disappear, but become restrictions that set the limits of permissible values.

The optimization parameters for this study are:

- time of completion of test tasks by users;
- the number of errors that occur during interaction with the system;
- SUS, CES, NPS interface evaluations;
- subjective assessment of the design on a 5-point scale;
- A/B testing results.

4. Selection of experiment participants

After determining the factors and feedback of the experiment, an important stage is the formation of a group of testers. According to the research of Jacob Nielsen, a leading usability expert, the optimal number of participants for high-quality testing is 15 people. This number allows you to identify up to 95% of usability problems while maintaining the effectiveness of the testing process, up to all the specified tests. It is important to note that according to Nielsen's research, after the 5th tester, the main significant interface problems are identified, and further testing serves to confirm and clarify the results [9].

Fig. 2 shows an example of a graph from these studies, which shows the percentage of usability problems found in a web application relative to the number of participants required to find them.

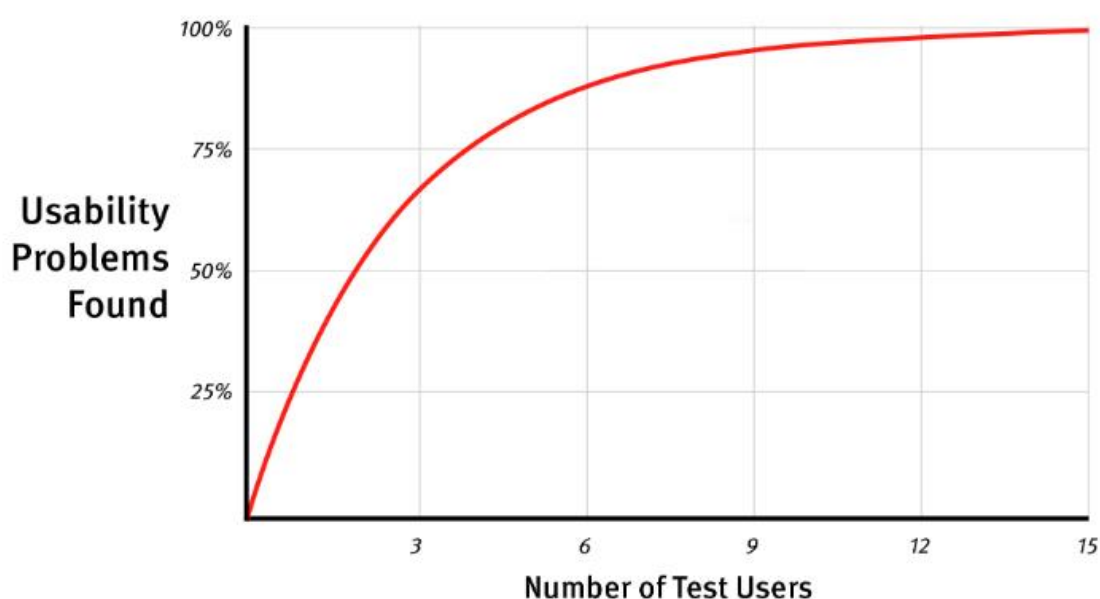


Fig. 2. Graph of optimal sample size for research into web application usability issues

For the research that will be conducted as part of the thesis, the following criteria will be applied to the formation of a group of test participants:

1. Staff: 15 people.

Participants are divided into three categories based on experience level: 5 people with limited experience working with web applications, 5 intermediate users who occasionally use such systems, and 5 professionals who have extensive experience in this field and interact with various web applications on a daily basis.

2. Demographic characteristics.

Participants were selected based on age range (20–55 years), gender balance (8 men, 7 women), and professional distribution (IT professionals, creative professions, humanities). This selection ensures the representativeness of the results for different categories of the target audience.

3. Participant selection criteria.

The experiment involved individuals with basic computer skills, but with varying levels of familiarity with the presented web application interfaces. Those who had participated in similar studies within the past six months were excluded to minimize outside influences on the test results.

5. Methodology for processing experimental results

Statistical processing. In the process of experimental research of UI/UX design of web applications, numerical values of key metrics will be obtained for each tested combination of tools and methodologies. Since the results of user interaction with interfaces always contain an element of randomness caused by individual characteristics of participants and external factors, it is necessary to apply statistical methods to validate the obtained data.

Using comparative analysis of average values using Student's t-test for independent samples will allow us to determine whether there are statistically significant differences between different interface options in terms of such indicators as task completion time and number of errors. In mathematical formalization, this is expressed through the null and alternative hypotheses [10].

Statistical processing of experimental data is aimed at analyzing the reliability of the obtained results, calculating their statistical indicators and determining whether they belong to the same general population. Only after such analysis can conclusions be drawn about the suitability of the data for further application. In addition, this process allows you to obtain additional information about the phenomenon or process under study.

Data preprocessing. At the stage of data pre-processing, a set of measures aimed at preparing information for further analysis is envisaged. The peculiarity of working with UX research data is the need to process both quantitative and qualitative indicators, which requires an individual approach to each type of data.

Primary processing begins with verification of the collected information for completeness and integrity. For questionnaire data, special attention is paid to identifying missing values, which may arise due to technical errors or respondents' reluctance to answer individual questions. In case of significant omissions, either an additional survey is provided or such observations are excluded from further analysis. When identifying and processing outliers – extreme values that may significantly affect the results of the study, for quantitative indicators (task completion time, number of errors) the method using Fisher's z-test is used. Qualitative assessments (SUS, subjective assessments) are analyzed for internal consistency using the coefficient of variation.

To ensure data comparability, a procedure for normalizing indicators is provided, especially for metrics measured in different scales. Time indicators are brought to a single scale, and subjective assessments are standardized taking into account the peculiarities of assessment scales. The result of pre-processing is a structured data set, ready for the application of statistical analysis methods.

Histograms. Since data is often random in nature, it is important to analyze its distribution. This determines the choice of the right processing methods, as different types of distributions require different approaches to analysis.

As part of the study, histograms will be used to analyze three main groups of indicators:

1. For task completion times by users, a histogram will allow you to identify typical values of time spent, as well as extreme cases that may indicate problem areas of the interface;
2. Histograms of SUS ratings will help to assess the overall level of user acceptance of the system. A normal distribution of these ratings will indicate stability of the results, while the opposite distribution may indicate differences in opinions among different groups of users;
3. For subjective design evaluations, histograms will reveal user preferences for visual solutions. An important analysis will be to compare histograms for different interface options.

Correlation. Correlation reflects a statistical relationship between two or more variables, where a change in one variable is systematically associated with a change in another. In the context of UX research, one might study the relationship between

framework type and user experience ratings. For quantitative data, the Pearson correlation coefficient is used, and for ordinal data, the Spearman correlation coefficient is used.

Correlation reflects the degree of dependence between random variables. It can be:

- positive, where the coefficient is close to +1. An increase in the values of one variable is accompanied by an increase in the other;
- negative, the coefficient is close to -1. An increase in one variable leads to a decrease in the other;
- absent, where the coefficient is 0. There is no statistically significant relationship between the variables.

The Pearson correlation coefficient is the most commonly used tool for examining linear relationships between variables in statistical research. It is used for variables measured on an interval or ratio scale and allows us to quantify both the strength and direction of the relationship between them.

The range of values for the Pearson correlation coefficient is limited to the interval from -1 to $+1$. The maximum positive value of $+1$ indicates a perfect direct linear relationship, when all data points lie exactly on the increasing line. On the contrary, a value of -1 indicates a perfect inverse relationship, where an increase in one variable is accompanied by a proportional decrease in the other. A zero value of the coefficient means the absence of any linear relationship between the studied indicators.

The graphical interpretation of the correlation coefficient is based on the analysis of the scattering of points on the coordinate plane. The closer the points are to an imaginary straight line that reflects the average relationship between the variables, the higher the absolute value of the correlation coefficient (see fig. 3).

Spearman's rank correlation method is an effective tool for detecting statistical relationships between variables, especially when the data are ordinal or do not meet the assumptions of parametric methods. The essence of this approach is to analyze the relationship between the ranks of values, rather than their absolute values.

Spearman's correlation coefficient belongs to nonparametric analysis tools and does not require strict assumptions about the nature of the distribution of the original data.

Like Pearson's coefficient, Spearman's coefficient varies from -1 to $+1$, with extreme values indicating a complete direct or inverse relationship between ranks, and a value of zero indicating the absence of any monotonic relationship. Positive values of the coefficient indicate that an increase in one variable

is accompanied by an increase in the other, while negative values indicate the opposite trend.

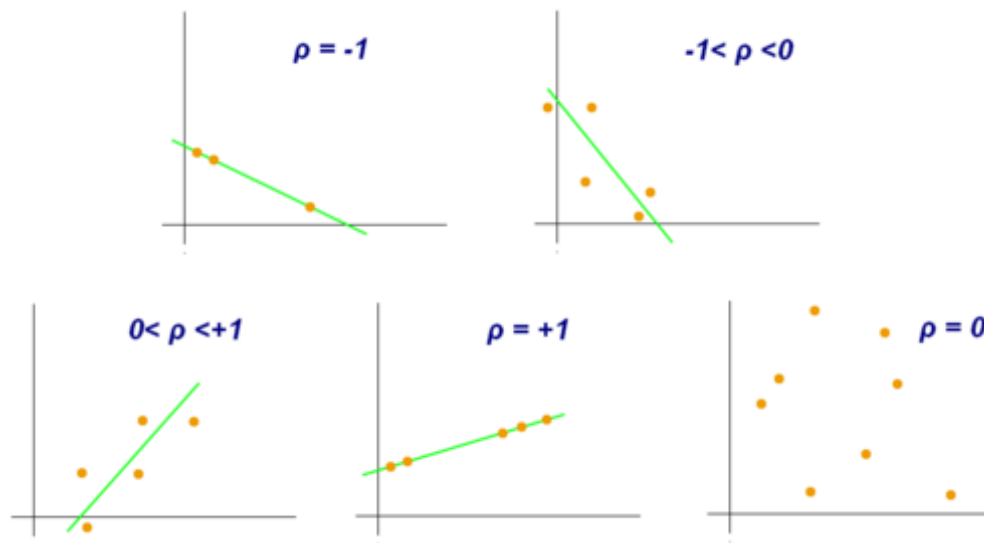


Fig. 3. Pearson correlation coefficient

The choice between Pearson and Spearman coefficients in this study will be made after a careful analysis of the nature of the distribution of the obtained data. For metrics measured in interval scales or ratios with a normal distribution, the parametric Pearson method will be used. In cases where the data are ordinal or significantly deviate from a normal distribution, priority will be given to the non-parametric Spearman coefficient.

Analysis of using Figma for UI prototyping. Figma is the leading UI prototyping tool. This section will analyze the use of Figma for UI prototyping, including the time to create mockups, ease of teamwork, and the efficiency of communicating designs to developers. It will be compared to Adobe XD, which is also a popular tool in the industry.

The main difference between Adobe XD and Figma is that it specializes in creating dynamic interfaces with advanced animation capabilities. The most significant advantage of this tool is the Auto-Animate feature, which allows you to design complex page transitions with minimal effort.

An important feature of Adobe XD is its close interaction with other Adobe Creative Cloud products, including Photoshop and Illustrator. This integration provides a seamless workflow for designers who use these applications to prepare graphic elements that can be easily imported into Adobe XD. This interconnection allows you to effectively combine work with vector graphics and prototyping in a single environment.

Let's look at Figma and Adobe XD in terms of their effectiveness for prototyping UI for web applications. The focus is on Figma, as it is one of the most popular tools among UI/UX designers due to its cloud-based platform, and it has a number of features that significantly impact the UI creation process. Its architecture is focused on collaborative editing, which allows not only designers but also developers, analysts, and customers to be involved in the process. Figma's prototyping tools include built-in interactive transitions between screens, the ability to implement interaction logic through «Smart Animate», and the ability to create components and variants, which allows you to design dynamic interfaces with adaptive behavior. Figma does not require installing applications, works in the browser, and automatically saves all changes.

For a comparative analysis of the effectiveness of using Figma and Adobe XD, Tbl. 1 was formed according to the main criteria necessary for UI prototyping. The comparison includes platform availability, teamwork capabilities, prototyping functions, integration with third-party tools, performance on large projects, and accessibility in terms of cost. This allows for an objective evaluation of the advantages and limitations of each tool in real design workflows.

Table 1

Comparison of Figma and Adobe XD in the context of UI prototyping

Criterion	Figma	Adobe XD
Platform	Web interface, cross-platform (Windows, macOS, Linux via browser).	Desktop application for Windows and macOS, no browser version.
Teamwork	Real-time collaborative editing, commenting, version history.	Commenting via Creative Cloud, collaborative editing of the project (Coediting mode).
Prototyping	Smart Animate, interactive transitions, various tool options.	Auto-Animate, 3D Transforms, Repeat Grid support for voice interactions (Voice mode).
Integration	Plugins, API, support for Zeplin, Jira, Slack.	Deep integration with Adobe Creative Cloud (Illustrator, Photoshop).
Work offline	An internet connection is required for most features.	Full offline work.
Productivity with large projects	High with stable internet, but requires browser restrictions.	Higher stability thanks to local storage.
Prototype creation speed	2.7 hours (average time for a 7-screen layout)	3.2 hours (average time for a 7-screen layout)
Accessibility	Free plan with limitations; paid version from \$15/month.	Available as part of Adobe Creative Cloud, starting at \$28.8/month; free version with limited functionality.

Fig. 4 shows the Figma interface when creating an interactive prototype, showing how the Smart Animate feature is used to create smooth transitions between screens.

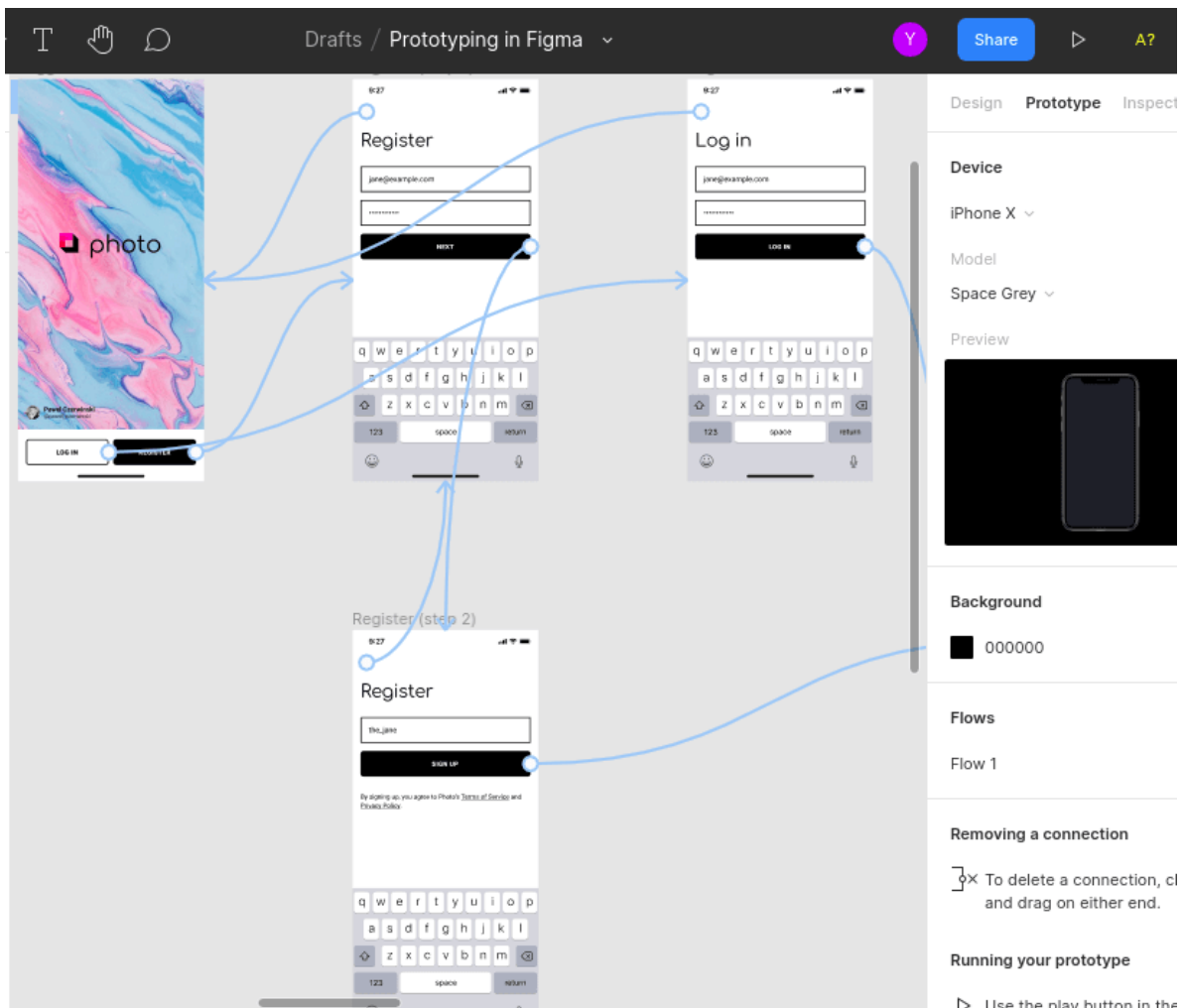


Fig. 4. Figma interface when creating a prototype using Smart Animate

Summarizing the results of the analysis, we can conclude that Figma is the optimal tool for creating web application prototypes. Its capabilities, including the layout of interfaces based on reusable elements, support for Smart Animate for smooth interactions, as well as flexible organization of collaboration, allow you to achieve high speed and accuracy during UI design, in particular, it is 17% faster (based on our own research) and easier to develop prototypes than in Adobe XD. In addition, Figma's browser-based accessibility and low entry threshold make it convenient for both team use and individual development. This makes Figma superior to Adobe XD in most practical prototyping cases and proves to be the most convenient environment for developing modern web interfaces.

Analysis of the use of CSS frameworks and UI libraries for interface implementation. Developing a front-end application without ready-made libraries allows you to create an interface according to your own rules from scratch. However, using such libraries not only saves development time, but also helps create clean and symmetrical layouts without additional effort. They cover a wide range of elements such as: buttons, forms, drop-down menus, modal windows, as well as tools for grid systems, styling, and user event handling.

In the modern web development ecosystem, there are many UI libraries, each with its own features, advantages, and limitations. In this section, we will perform a comparative analysis of selected CSS frameworks and UI libraries: Tailwind CSS, Bootstrap, Material UI, for implementing user interfaces.

Tailwind CSS is a relatively new front-end framework, first released on November 1, 2017, with its first version being created by Adam Watan and Steve Shoger. The library has since garnered significant attention and has even been adopted by major companies including Algolia and Mozilla.

Bootstrap was originally developed by Mark Otto and Jacob Thornton while they were both working at Twitter. The first version was first released on Friday, August 19, 2011. Bootstrap originally used CSS and optionally jQuery design templates for its components. However, in the latest version, Bootstrap 5, jQuery support was dropped in favor of plain JavaScript. It provides a set of pre-built components and styles that developers can use to quickly build web applications, and it also uses a grid system to improve layout and has a consistent visual style across all of its components.

Material UI is Google's implementation of Material Design as a set of React components that is primarily focused on mobile devices – we write the code for mobile devices first, and then scale the components as needed using CSS media queries. To ensure proper rendering and touch scaling on all devices, a «responsive viewport» meta tag is added to the <head> element.

Unlike traditional CSS frameworks like Bootstrap, which provide pre-built styles and components, Tailwind CSS instead provides a set of low-level utility classes that allow developers to create their own designs right out of the box. Tailwind classes can be combined and composed to create any type of interface, from simple buttons to complex components. This approach is inspired by the «atomic» CSS ideology, which involves using many small, single-purpose CSS classes to style elements quickly and intuitively.

Tailwind CSS works by scanning class names across all HTML files, JavaScript components, and other templates, then generating the appropriate

styles and writing them to a static CSS file. Below is an example of how to create a basic blue button in Bootstrap, MUI, and Tailwind CSS. This comparison demonstrates the different approaches to styling components: Bootstrap provides predefined utility classes, MUI offers ready-to-use React components with built-in customization, while Tailwind CSS gives developers low-level control over design through utility-first classes.

Implementation using Tailwind CSS:

```
<button class=«bg-blue-500 hover:bg-blue-600 text-white font-bold py-2 px-4 rounded»> Normal button </button>
```

Example of pure Bootstrap:

```
<button class=“btn btn-primary”> Normal button </button>
```

Example in Material UI:

```
<Button variant=«contained» color=«primary»> Normal button </Button>
```

From this example, you can see that unlike other libraries, Tailwind allows you to directly customize elements using helper classes. The developer can manually specify the color variant, the background color on hover, etc. It is also possible to use one library together with another, although this is not recommended, because different libraries and frameworks have the same class names, for example, the classes «container», «px» and «mx» exist in both Bootstrap and Tailwind, the use of which can cause conflicts during design. Tbl. 2 shows a comparative analysis of these libraries.

Table 2

Comparative analysis of Bootstrap, Tailwind CSS and Material UI

Criterion	Bootstrap (v5.X)	Tailwind CSS (v4.X)	Material UI (v7.X)
Library size (minified CSS)	~165+ kB	~190+ kB	~335+ kB
Standard interface development time	~3.2 hours	~4.1 hours	~3.8 hours
Customization flexibility	Limited to topics	Full customization	Customization possible via JSS, more complex configuration
Performance (Lighthouse Score)	82/100	95/100	78/100
Popularity (GitHub Stars)	172k	88k	95k
Support for adaptability	Built-in adaptive grid	Utilitarian approach	Component library based on Material Design
React integration	Through the React-Bootstrap library	Direct, via JavaScript XML	Native (all components are designed specifically for React)
Number of templates/themes	5000+	A set of built-in utility classes	300+ (MUI Store)

Thus, each of the analyzed libraries has its strengths and weaknesses. Bootstrap is the easiest to use, thanks to a wide range of ready-to-use components. Tailwind CSS, on the contrary, requires more experience and setup time, as it works at the atomic class level, but provides a unique approach to web design, and more opportunities for styling. MUI is ideal for React applications and provides a clean, modern design that meets Google's recommendations for Material Design. The use of these libraries allows you to effectively implement modern web interfaces taking into account the criteria of usability and visual appeal.

6. Tools for collecting UX feedback and conducting usability testing

One of the main stages of the experiment is collecting UX feedback and conducting usability testing with subsequent analysis of the data obtained. The quality of the collected information, and therefore the validity of the research conclusions, depends on the correct choice of tools. This section discusses in detail the following tools – the Lyssna platform (formerly UsabilityHub) and Google Forms, which were selected for conducting UX research.

Lyssna is a platform for remote user research, where you can conduct both moderated and unmoderated research in the form of interviews, usability tests, design surveys, etc. The platform provides a wide range of testing methods, each of which is focused on studying certain aspects of UX, fig. 5 shows all the available tests of this platform:

1. Prototype test – allows you to upload interactive prototypes from Figma, challenge users, and analyze their actions. This method is especially useful for evaluating System Usability Scale (SUS), Net Promoter Score (NPS), and heat maps, as it allows you to customize metrics to suit your research needs.
2. First click test – determines whether users find key interface elements the first time.
3. Five-second test – helps to evaluate the first impression of the design.
4. Design survey – allows you to ask questions while viewing an image or video.
5. Preference test – used for A/B testing when users choose between two design options.
6. Card sort & Tree test – examines the information architecture of the product.
7. Live website test – records user actions on a real website.

For unmoderated tests, Lyssna supports PNG, JPG, GIF, MP3, MP4, CSV, Figma prototype files, as well as live website testing. This means you can easily test designs exported from tools like Figma, Sketch, or Invision, screenshots from live

sites, live websites themselves, or even photos of sketches. For moderated studies, the features are a bit more extensive – you can share your screen to get feedback on your designs regardless of the file type, or link participants to your website so you can watch them share their screen and complete a task or set of tasks.

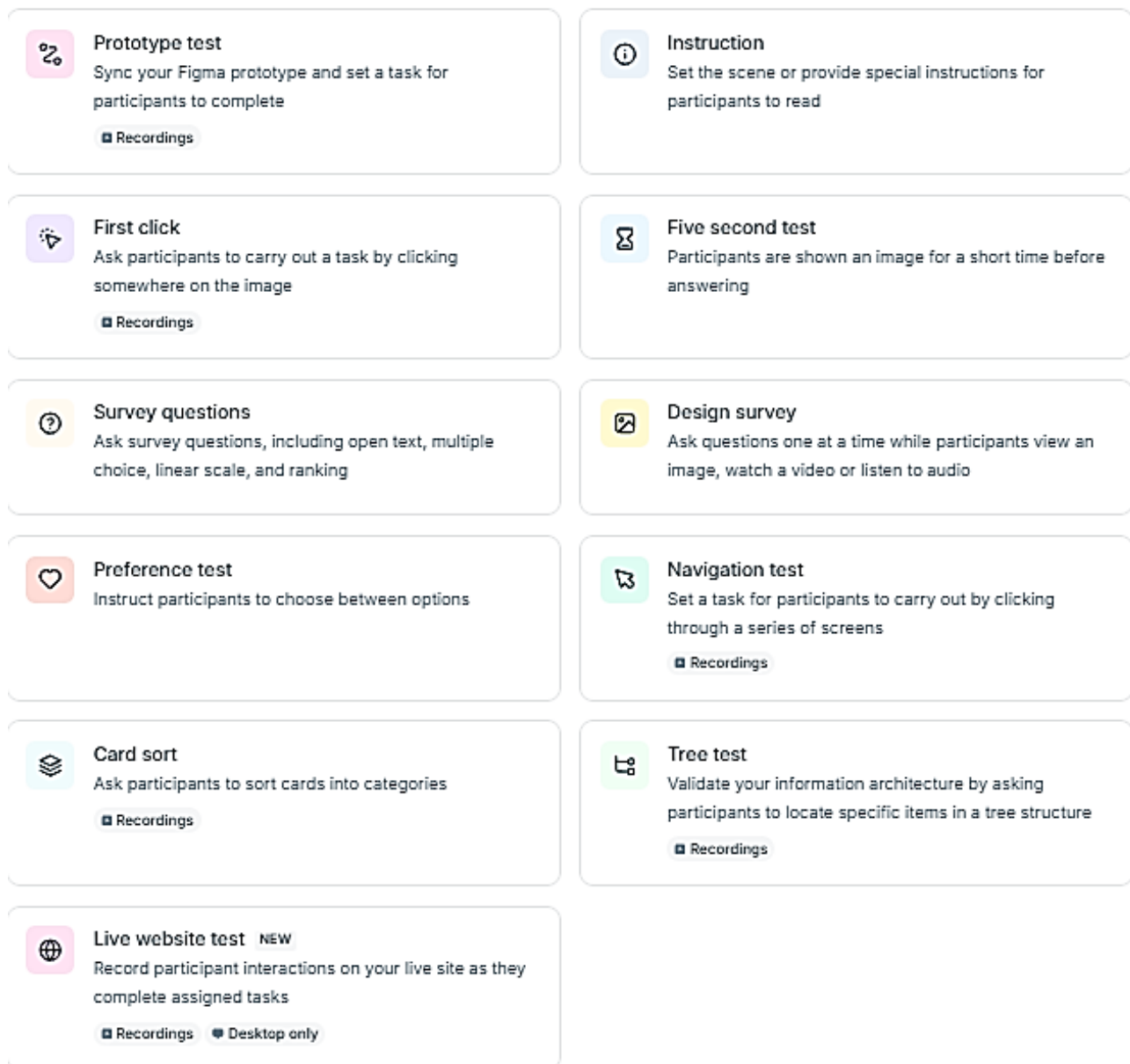


Fig. 5. Available unmoderated tests of the Lyssna platform

The study will pay particular attention to Prototype Test, as this type of testing provides a variety of opportunities for creating individual user interaction scenarios with a prototype built in Figma. This allows you to integrate an already developed interface from the Figma environment directly into the Lyssna platform, after which you can set specific tasks for the user, define the logic of the scenarios, select evaluation metrics – such as System Usability Scale (SUS), Net Promoter

Score (NPS), click heat maps, and others. The process of connecting and configuring such testing is demonstrated in fig. 7. Another important advantage is the automatic generation of reporting on the test results, which includes both quantitative and qualitative indicators. For comparative analysis of different design solutions, Lyssna implemented Preference Test, which performs the functions of A/B testing and was chosen as the second tool in this study.

Another key functionality of Lyssna, presented in fig. 6, is the ability to customize the distribution of testing. The researcher can generate a unique link to the test and send it to the target audience, or use the service of finding respondents directly on the platform. At the same time, the available targeting tools allow you to specify the specification by demographic parameters, professional skills, employment status, etc. This is especially important in cases where the product is aimed at a narrow segment of users, or when an analysis of a large group of potential users is required in a short time. In addition, Lyssna provides real-time statistics on the progress of responses, which simplifies monitoring and accelerates data collection. The flexibility of respondent management also helps reduce testing costs, while maintaining high reliability of the collected data.

The second tool chosen for collecting UX feedback was Google Forms. Although Google Forms is not a specialized tool for UX research, it was chosen as a supporting tool for conducting A/B testing and collecting subjective user ratings and UX feedback.

Google Forms is a cloud-based software that allows you to create, conduct, and analyze surveys, tests, questionnaires, feedback forms, and other methods of collecting data online. This tool is part of the Google Workspace suite and is available through a web browser. The universal and free platform provides an interface for constructing questionnaires using different types of questions (open, closed, scale, rank), logical branching, and automated statistical processing functions (visualization in the form of graphs, tables, and charts).

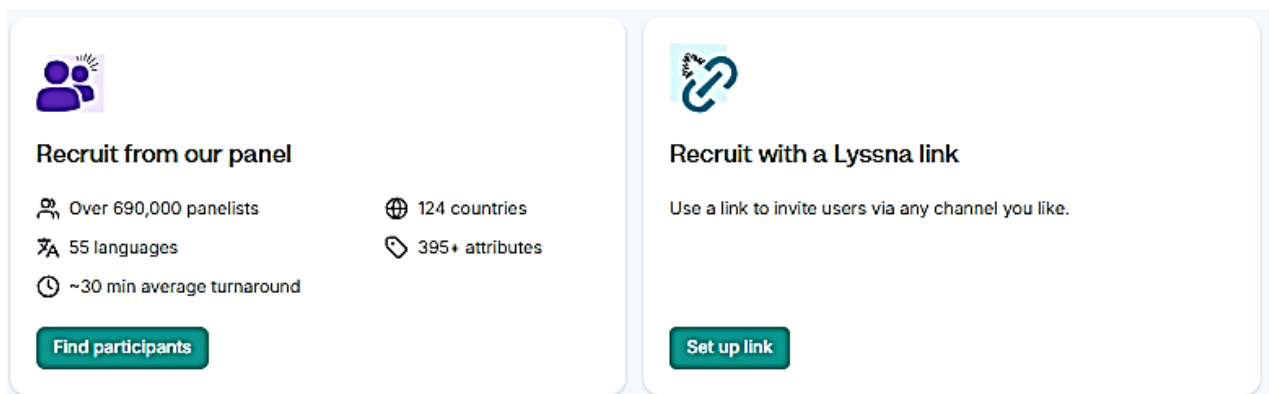


Fig. 6. Available Lyssna features for testing distribution

1. Prototype test

☐ Task flow
Ask participants to complete a task

☒ Free flow
Ask questions while participants roam around

Instructions (optional) ⓘ

italic
****bold****
Position

Bottom right ▾

☒ Customize system instructions ⓘ

- Це інтерактивний прототип.
- Клацайте або торкайтеся екранів, щоб показати, як би ви виконали завдання.
- Не хвилюйтеся, якщо ви натискаєте на щось, а воно не реагує — це просто означає, що його не можна натиснути.

MEDDICAL - Hospital website template (Community) [↗](#)

● Last synced on 28 мая 2025 р. в 03:45 PM

Scaling

Fit – Scale down to fit ▾

Starting point
Home

Follow-up questions

Questions for participants to answer before proceeding to the next task.

1.1. Linear scale question

Question

Linear scale ▾
Required ☒

Start value

1 ▾

Зовсім не згоден

End value

5 ▾

Повністю згоден

1.2. Linear scale question

Question

Linear scale ▾
Required ☐

Fig. 7. Example of configuring Prototype test to SUS metric

Within the framework of this study, Google Forms is used to implement basic A/B tests, where respondents are shown two different implementations of the interface (for example, a login page or a home page) and asked questions about convenience, visual perception, ease of navigation, etc. Google Forms will also

66

be used to collect feedback from respondents who have passed the tests from the Lyssna platform, on open-ended questions, as well as to implement a survey on the SUS scale in case of access restrictions to Lyssna. They can leave their wishes for adding or removing various features to the future web application, based on its prototype, and give an expert assessment of its visual component.

Thus, using the above tools provides the opportunity not only to compare different interface options, but also to identify specific problem areas, understand the reasons for certain user behavior, and, as a result, make informed decisions about product improvement.

7. Algorithm for conducting the experiment

To conduct a comparative experiment on the study of UI/UX design of web applications, each stage of the study is performed taking into account the goals of the experiment, the types of interfaces being analyzed, and the specifics of the tools used to implement and evaluate the design. This section provides a step-by-step description of the algorithm that will be used within the study to systematically test the hypotheses formed and collect data for further analytical processing. Experimental research begins with a clear definition of the scientific hypotheses to be tested. Three key hypotheses were formulated within the framework of this work:

- H1. Using Figma during the prototyping stage reduces the number of implementation errors compared to Adobe XD by 15%;
- H2. Using Tailwind CSS allows you to speed up the implementation of responsive UI compared to Bootstrap;
- H3. Interfaces created according to the principles of Material Design (Material UI) demonstrate 15–20% better usability indicators on the SUS scale compared to interfaces without clear design guidelines.

In order to compare the effectiveness of UI/UX approaches, it is necessary to create several interface mockups using modern prototyping tools – Figma and Adobe XD. The mockups reflect the same structure and functionality, but are implemented in different styles, according to the capabilities of the specified tools. Particular attention is paid to adaptability, speed of implementation, navigation logic, as well as compliance with modern design principles.

The next step is to implement each of the layouts as web interfaces using the three most popular CSS frameworks: Tailwind CSS, Bootstrap, and Material UI. The choice of these frameworks is due to their widespread use in web development and different concepts of organizing styles. Tailwind provides a high level of

customization, Bootstrap provides rapid development using ready-made components, and Material UI implements Google's design principles.

Creating user test scenarios. For each implementation of the interface, you need to create a separate test scenario that involves the user performing a series of tasks sequentially. The scenarios are focused on typical actions that reflect the main goals of using the interface.

Example of a scenario in the Lyssna platform (Prototype Test): «Imagine that you are a new user of the web application. Your goal is to register on the platform, get acquainted with the main functions of your personal account and change your profile settings. Find the «Register» button, go to the registration form, create a test profile, open the «My Profile» section and change your username».

This scenario allows you to test navigation logic, intuitive element placement, task execution time, and identify potential errors.

After completing the tasks, the respondent will be offered:

- assess the difficulty of each step on a scale from 1 to 5;
- fill out the standard SUS questionnaire;
- indicate the probability of recommending the service (NPS);
- provide free comments about the interaction experience.

The next stage will be UX testing. At this stage, a series of UX tests are launched using the Lyssna platform, using the following test types: Prototype Test (interactive prototype testing), Preference Test (choosing the best option), SUS (System Usability Scale), NPS (Net Promoter Score), as well as heatmaps. Each test is configured according to the interaction scenarios described above. To conduct the tests, you need to recruit 15 users with different experience (from beginners to professionals) with a balance by gender and age (20–55 years). The collected data will be stored in the platform's analytics system for further analysis.

The test results are exported for quantitative and qualitative analysis. Indicators such as average task completion time, number of clicks, accuracy of execution, SUS/NPS scores, visual patterns on heat maps, etc. are evaluated. For quantitative data, the Pearson correlation coefficient is used, and for ordinal data, the Spearman correlation coefficient is used. For analysis, a tabular form, correlation analysis, construction of diagrams and comparative graphs are used. The use of descriptive statistics and, if necessary, methods of hypothetical testing (t-test) is also envisaged.

The final stage, based on the results obtained, will be a comparative analysis of the effectiveness of the tools used (Tailwind, Bootstrap, Material UI), prototyping tools (Figma, Adobe XD), as well as UX assessment methods. In particular, the advantages and limitations of each approach, compliance with

user expectations, as well as recommendations and conclusions regarding the optimal choice of UI/UX design tools for the implementation of web applications in future projects will be determined.

Conclusions

As a result of the second section, a complete experimental model was developed to study the effectiveness of methods, tools, and technologies for developing UI/UX design for web applications. Relevant factors and experimental feedback were identified, including the use of various CSS frameworks (Tailwind CSS, Bootstrap, Material UI), prototyping tools (Figma, Adobe XD), and UX evaluation metrics (SUS, NPS, click heatmaps).

A plan for conducting an experimental study was developed and described, which includes setting goals, selecting participants (15 respondents with different levels of experience with web applications) and a method for processing results using statistical, correlation analysis, and histograms. An analysis of prototyping tools (Figma, Adobe XD) and CSS frameworks (Bootstrap, Tailwind CSS, Material UI) was conducted, their advantages and disadvantages were identified. The use of the Lyssna and Google Forms platforms for implementing UX testing was justified, and the methodology for creating test scenarios, integrating with Figma, and implementing evaluation using SUS, NPS, and heat map metrics was described.

In the future, the results of this study can be expanded towards integration with automated UX analysis methods based on processing large amounts of data (datasets) and the use of adaptive machine learning models (AutoML). Such integration will allow more accurate prediction of user behavior patterns and reduce the time required for manual analysis. Additionally, the application of predictive analytics may uncover hidden correlations between interface design parameters and user satisfaction levels.

References

1. Yablonski J. Laws of UX: Design Principles for Persuasive and Ethical Products / Jon Yablonski., 2020. – 150 p. – (O'Reilly Media; 1st edition). – (978-1492055310).
2. Atomic Design [Electronic resource] / Brad Frost. – 2013. – Resource access mode: <https://bradfrost.com/blog/post/atomic-web-design/>
3. Design Thinking (DT) [Electronic resource] // Interaction Design Foundation. – 2024. – Resource access mode: <https://www.interaction-design.org/literature/topics/design-thinking?srltid=AfmBOooxImwpqvabHyTTjq1mp2RmsaVFsrNOVnY1-IFHeW49V2yrfMWy>

4. Human-Centered Design (HCD) [Electronic resource] / Don Norman // Interaction Design Foundation. – 2024. – Resource access mode:
https://www.interaction-design.org/literature/topics/human-centered-design?srsId=AfmBOoqOIG0MtvKw22LwJMavKrhHOnwoL11v1ZHOxqC7Ni5JqThXd_29
5. Cooper A. The Inmates Are Running the Asylum / Alan Cooper., 2004. – 288 p. – (Sams – Pearson Education). – (9780672326141).
6. Krug S. Don't Make Me Think, Revisited / Steve Krug., 2013. – 216 p. – (New Riders; 3rd edition). – (978-0321965516).
7. Norman D. The Psychology of Everyday Things: Revised and Expanded Edition / Donald Norman., 2013. – 368 p. – (Basic Books). – (978-0465050659).
8. Sequence of experiment organization [Electronic resource] – Resource access mode:
<http://topknowledge.ua/nauchnye-issledovaniya/3041-posledovatelnost-organizatsii-eksperimenta.html>
9. Why You Only Need to Test with 5 Users [Electronic resource] / J. Nielsen. – Resource access mode:
<https://www.nngroup.com/articles/why-you-only-need-to-test-with-5-users/>
10. Student's t-test [Electronic resource] – Resource access mode:
https://en.wikipedia.org/wiki/Student%27s_t-test

INTEGRATING EPIDEMIC SIMULATION INTO CRISIS MANAGEMENT PROGRAMS: LESSONS FROM UKRAINE

Chumachenko D., Meniailov Ie., Bazilevych K., Parfeniuk Yu., Chumachenko T.

Crisis-management programs (PHEOC/ICS) need decision-ready analytics under uncertainty. Russian war in Ukraine, marked by attacks on health care, sustained displacement, and environmental shocks, exposes the limits of ad-hoc modelling and the value of routine, integrated simulation support. To synthesize recent evidence on how epidemic simulation can be embedded in crisis-management workflows and to distill practical lessons from the Ukraine context. Methods: Rapid narrative review of 2023–2025 peer-reviewed studies and authoritative guidance. Findings were thematically mapped to PHEOC/ICS functions. The literature converges on four effective integration points: planning, operations, logistics, and risk communication. Multi-model ensembles are most reliable for short horizons; scenario exercises support medium-term strategy when assumptions are explicit and versioned. Displacement-aware inputs can triangulate shifting catchments; modular focus areas include measles catch-up, polio immunity maintenance, TB care cascades, and WASH-sensitive risks after infrastructure damage. Enablers are a fixed cadence, standard one-page briefs, clear roles (modelling liaison), equity checkpoints, and decision logs. Common barriers are data plumbing, staffing, and product–process mismatches. We provide an operational template and a concise checklist that translate current evidence into routine command practices: weekly ensemble briefs, scheduled scenario cycles, displacement-aware data pipelines, and explicit mapping of outputs to ICS artefacts. These steps are immediately actionable in Ukraine and transferable to other high-constraint settings.

Introduction

Public health emergencies place decision makers under tight time and information constraints. Incident management frameworks, such as the Public Health Emergency Operations Centre (PHEOC) and the Incident Command System (ICS), were built to structure those decisions by coordinating information, operations, and logistics across agencies [1]. Epidemic simulation can add value in this setting by translating surveillance and capacity data into short-term forecasts and scenario comparisons aligned with incident action planning [2]. Evidence from the COVID-19 period shows the promise of modelling and the limits that arise when assumptions shift or uncertainty is not communicated well [3]. Recent reviews argue that models inform policy best when integrated into established response workflows, with clear products, cadence, and roles [4]. However, scoping work on PHEOCs finds that such integration remains uneven, with data management, staffing, and governance gaps [5]. These points motivate a focused synthesis on how to embed epidemic simulation into crisis management programs in feasible and useful ways.

Ukraine offers a critical case. Since the escalation of the Russian war in 2022, the health system has faced sustained attacks on facilities and personnel, persistent displacement, and strain on routine services. WHO has documented more than two thousand attacks on health care in Ukraine [6], and UNHCR reports millions of people displaced inside the country and across borders [7]. These conditions complicate surveillance, supply chains, and access to care, precisely the areas where crisis programs and decision support must perform.

The epidemiological risk landscape has also shifted. The WHO European Region reported the highest number of measles cases in over 25 years in 2024, underscoring the danger of immunity gaps in mobile and hard-to-reach populations [8]. At the same time, Ukraine's 2021 cVDPV2 polio outbreak was officially closed in September 2023, showing that targeted vaccination can succeed despite conflict [9]. Tuberculosis (including drug-resistant TB) remains a regional concern, with WHO's 2024 global report and country assessments noting elevated incidence and challenges linked to displacement and service disruption [10]. Conflict-related environmental shocks further raise risk: the 2023 destruction of the Kakhovka Dam degraded water and sanitation infrastructure and, according to recent studies, mobilized legacy heavy metals and other contaminants, hazards that can elevate water- and rodent-borne disease risks in affected communities [11]. Together, these signals illustrate the need and the opportunity to use simulation to anticipate burdens and compare interventions under constraints.

Operationalizing simulation in this setting requires reliable situational inputs. Recent work shows that displacement and population mixing, key drivers of infectious spread, can be monitored using complementary data sources [12]. Studies have nowcast internal displacement in Ukraine using social media activity, GPS-based human mobility, and satellite-based vehicle counts, offering different coverage and bias profiles [13]. Combined with routine surveillance and capacity data, these streams can help parameterize timely models to inform planning meetings and resource allocation.

This rapid review synthesizes practical lessons on integrating epidemic simulation into crisis management programs, drawing on Ukraine as the primary context. We map where models can support the ICS/PHEOC cycle, summarize enablers and barriers identified in recent guidance and reviews, and distil actionable steps suitable for short, repeatable decision cycles. The aim is to translate current evidence into simple, usable guidance for teams making high-stakes choices under uncertainty.

Methods

We conducted a rapid narrative review to synthesize how epidemic simulation is integrated into crisis management programs, with Ukraine as the focal context. Methods followed evidence-informed rapid-review guidance and recent handbooks for narrative synthesis. Specifically, we drew on the 2024 Cochrane Rapid Reviews Methods recommendations, the JBI Manual for Evidence Synthesis, and reporting resources for non-meta-analytic syntheses.

We included sources that examine epidemic modelling/simulation used for or proposed to support incident management, emergency operations, or policy decision making, or describe crisis management frameworks relevant to modelling integration or analyze Ukraine's public health emergency management since 2022, where modelling or decision support is discussed. We considered peer-reviewed articles and authoritative organizational reports/guidance to ensure operational relevance. Framework references for PHEOC/ICS/HICS were used to anchor concepts.

We prioritized items published January 1, 2023 – August 25, 2025 to reflect current practice. Seminal earlier framework documents were retained when no newer equivalents existed.

We searched PubMed/MEDLINE, Scopus, and Web of Science for peer-reviewed literature, and we searched targeted grey literature portals and organizational sites.

Given the narrative, mixed-evidence scope, we assessed reporting quality and credibility rather than applying study-design-specific risk-of-bias tools across all sources. We used the SANRA quality elements for narrative and guidance documents to guide appraisal and structure notes on strengths/limitations. For textual/qualitative evidence, we used JBI 2024 textual-evidence guidance to judge clarity, coherence, and source authority.

Results

The literature converges on a practical view of where simulation fits inside emergency programs: models are most useful when their outputs are delivered at the same rhythm and decision points as an IMS or PHEOC. Guidance from WHO's EOC-NET and CDC's IMS materials describes those decision points, situation assessment, incident action planning, resource tasking, and public briefings, and emphasizes the information flows and roles that turn analyses into orders [14, 15]. In this frame, simulation is not an external «report» but a recurring input to the incident action plan and the briefing cycle.

Recent evaluations have found that multi-model approaches are the most dependable way to supply these recurring inputs [16]. Studies of the U.S. Forecast Hub and Scenario Modelling Hub show that ensembles stabilize performance and yield better-calibrated short-term forecasts than most single models [17]. At the same time, longer-range scenario exercises support planning when assumptions are explicit and versioned. These hubs also document what increases use in practice: fixed delivery schedules, common file formats, and concise uncertainty statements that explain what could change the recommendation.

Several reviews add why this «hub mindset» aligns with decision-making. First, decisions are often comparative rather than absolute, making ensembles and structured scenarios a natural fit [18]. Second, the design of scenarios matters: framing choices with decision analysis principles improves relevance and transparency [19]. Third, the interface needs care, knowledge-translation roles, standard slide templates, and decision logs to reduce friction between modellers and incident managers [20].

Applied to Ukraine's conditions since 2022, the same integration logic becomes more urgent. WHO/Europe has verified more than 2,254 attacks on health care as of February 2025, disrupting facilities, staff safety, supply routes, and routine data systems, precisely the elements an operations section depends on [6]. In parallel, UNHCR reports that millions of people remain displaced inside the country, and millions more are abroad, creating shifting catchments and uncertain denominators for service planning [7]. In such a setting, routine model products that plug into the incident briefing are more valuable than one-off analytics.

Recent studies also expand the data layer for displacement-aware planning. High-frequency social-media advertising data have been used to nowcast internal displacement at the subnational scale during the first months of the invasion [12]. Very high-resolution satellite car counts and other remote-sensing approaches provide complementary signals [21]. Newer work triangulates across digital traces to track flows across borders [22]. Each source has bias and coverage limitations, but together, they offer timely proxies when administrative data lags, useful for prioritizing outreach, staging supplies, and adjusting micro-plans.

The communicable disease picture in and around Ukraine sharpens the operational need for this pairing of models and command processes. WHO/Europe and UNICEF recorded 127,350 measles cases in the Region in 2024, the highest since 1997, while ECDC reports sustained transmission into early 2025 [8, 23]. Those signals argue for displacement-aware catch-up vaccination scenarios linked to incident triggers and logistics checklists. Conversely, the cVDPV2 polio outbreak

detected in 2021 was formally closed in September 2023, a reminder that targeted campaigns can deliver results even under wartime constraints when surveillance, micro-planning, and operations align [9].

Tuberculosis remains a persistent concern that crosses borders and care systems. The WHO Global Tuberculosis Report 2024 and the joint ECDC/WHO regional surveillance report describe ongoing burdens of TB and MDR/RR-TB and highlight pressures on case finding and continuity of care, precisely the weak points during conflict and displacement [10, 24]. For integration, the most relevant model outputs are those that stress-test the care cascade and quantify the effect of treatment interruptions on outcomes and resourcing.

Environmental shocks compound these risks. Peer-reviewed analyses of the Kakhovka Dam breach document pulses of sediment and nutrients into the north-western Black Sea and raise plausible concerns about contaminant mobilization and WASH-related disease hazards downstream [25]. UNEP's independent assessment and scientific reporting underscore how cascading environmental effects alter the operating environment for public health programs [26]. In such contexts, the literature supports using scenario-based simulations to test whether safe-water distribution, vector control, and surveillance boosts are adequate under different access and supply constraints.

The same sources clearly state what helps and hinders integration. On the enabling side, PHEOC/IMS provides the standing products and cadence into which ensembles and scenarios can be dropped with minimal translation. Recent European and WHO guidance also elevates simulation exercises (SimEx) and after-action reviews (AARs) as the main tools to institutionalize the modeller-operator handshake [27]. Teams can rehearse thresholds, see whether forecasts actually change orders, and refine formats before the next activation.

On the barrier side, a 2023 scoping review finds recurrent gaps in data management, staffing, and information-sharing inside PHEOCs, along with product-process mismatches [5]. These findings explain why otherwise strong analyses have limited impact: without clear data plumbing and a place on the agenda, forecasts arrive too late or in the wrong shape. Recent work on exercise design offers practical remedies, clarifying objectives, standardizing artefacts, and embedding evaluation, so that integration is practised, not improvised [28].

Equity is another thread running through recent guidance that interacts directly with modelling. Hospital and system-level studies propose concrete ways to embed equity roles and metrics into HICS, from job action sheets to activation checklists [29]. These structures make it easier to translate model outputs about

coverage gaps among displaced or hard-to-reach groups into tasking [30]. This is especially relevant in Ukraine and host countries, where access barriers vary by location and status.

The extended evidence suggests a simple pattern. First, ensembles and structured scenarios, delivered on a fixed cadence and documented with clear assumptions, perform more reliably and are easier to use in command settings than bespoke single-model reports. Second, the fit to the decision matters: when outputs are linked explicitly to ICS/PHEOC products (incident objectives, resource checklists, job action sheets), uptake improves. Third, Ukraine's operating constraints, attacks on care, sustained displacement, and environmental disruption, shape both the parameters and the feasible options the models should test. The most useful simulations, therefore, combine displacement-aware catchment estimates, care-cascade stress tests, and WASH-sensitive scenarios, all embedded in the routine briefing rhythm of the response.

The literature is consistent on communication: decision makers respond better to probabilistic outputs with plain language caveats and a short statement of «what would change our advice», rather than to point estimates [31]. New guidance from WHO on communicating uncertainty in health emergencies and qualitative work with policy users reinforces this [32]. When these practices are institutionalized, a standard slide, a named liaison, and a decision log, simulation stops being an occasional input and becomes part of the program's operating system.

Table 1 presents a checklist that distils recent evidence on embedding epidemic simulation in crisis-management programs, focusing on conflict-affected settings such as Ukraine.

Table 1

Checklist

Action	Implementation	Why it matters
1	2	3
Anchor modelling inside the command structure	Place a modelling/analysis cell under Planning in the PHEOC/ICS org chart; pre-define activation levels and reporting lines.	PHEOCs are designed to coordinate analysis into action. Explicitly placing modelling in the workflow improves uptake.
Adopt a fixed cadence & standard products	Issue a one-page briefing on a set schedule (e.g., weekly forecasts; monthly scenarios) with assumptions, uncertainty bands, and «what would change the advice».	Forecast/scenario hubs show higher policy use when outputs are routinized and standardized.

Continuation of the Table 1

1	2	3
Prefer ensembles & hubs for operations	Use an ensemble for near-term indicators; use multi-team scenario hubs for horizon scanning.	Ensembles are better calibrated on average; scenario hubs support comparative planning when assumptions are explicit.
Map outputs directly to ICS artifacts	Tie thresholds to ICS-202 (Incident Objectives); turn options into ICS-204 (Assignments); log rationale in ICS-214 (Activity Log).	Direct mapping turns numbers into orders without reformatting.
Specify the minimum viable data feeds	Maintain live feeds for: surveillance/labs; immunization & stocks; facility capacity; displacement/mobility; WASH/env.	Guidance identifies these as core information streams for effective PHEOC operations.
Triangulate displacement & access	When registries lag, combine Meta ads, GPS/mobile mobility, and satellite car counts to refine catchments and outreach.	Recent Ukraine studies validate each proxy and their complementarity for subnational nowcasting.
Stress-test WASH-sensitive scenarios	Include safe-water, surveillance, and vector-control options under different access/supply constraints.	UNEP rapid assessment and peer-reviewed analyses document contamination risks and hydrologic disruption after the dam breach.
Build an equity checkpoint into every cycle	Name an Equity Lead in HICS/PHEOC; add an «equity impact» line to each modelling brief and ICS-202.	Empirical and review work shows practical ways to embed equity roles and metrics into HICS.
Use plain-language uncertainty	Report probabilistic ranges; state key assumptions and explicit triggers for revision; align with RCCE principles.	Clear uncertainty communication increases trust and actionability in emergencies.
Exercise the interface; learn from it	Run SimEx that rehearse how forecasts change orders; conduct AARs to check timeliness, comprehension, and impact.	WHO/EU guidance and recent SimEx reports highlight SimEx/AARs as key institutional levers.
Record decisions & rationale	Keep a decision log linked to the modelling brief and ICS-214.	Standardized activity logs support accountability and rapid updates.
Plan data plumbing & governance	Document owners, sharing agreements, refresh rates, and fallbacks; resource the information function.	PHEOC assessments cite data management and staffing as recurrent gaps.
Define roles & liaisons	Assign a Modelling Liaison to Planning/Incident Commander; pair with Risk-Comms lead for joint briefings.	Formal role definitions reduce friction and align products with decisions.
Match questions to methods	Use ensembles for staffing/supply in the next 1–4 weeks; scenarios to compare strategic options; nowcasting for situational awareness.	Different decision horizons require different tools; CDC articulates the scenario role, and recent ensemble work supports near-term use.
Close the loop with evaluation	Add a standing «Was this useful?» item to planning meetings; track actions linked to modelling in AAR.	Continuous evaluation is embedded in current preparedness and exercise guidance.

Discussion

Our results show that simulation becomes useful when it is delivered at the same points where command teams make choices, during planning, operations, logistics, and briefings, and on a regular cadence. This placement is not incidental: PHEOC/IMS guidance is built around recurring products (situation assessments, incident objectives, resource requests) and a meeting rhythm that converts information into tasking. Aligning model outputs to that rhythm lowers the cost of use. It explains why programs that treat simulation as a standing input, rather than an occasional report, see more consistent uptake. Recent WHO materials on EOC practice, and national handbooks derived from them, reinforce this operational logic [14].

The evidence is strongest for two model types serving different horizons. First, near-term ensembles offer more stable accuracy than individual models for the one-to-four-week window that operations sections care about [33]. Evaluations of the US Forecast Hub repeatedly show better calibration and fewer large errors when teams are combined [34]. Second, scenario hubs can inform planning several months out when the big drivers are stated upfront and versioned, a point underscored by the formal assessment of the US Scenario Modelling Hub [3]. These findings support a dual track in crisis programs: ensembles for staffing, beds, and supplies, scenarios for strategy and contingency planning.

The Ukraine case adds concrete constraints that shape how those tools should be used. Sustained attacks on health care have reduced facility availability and disrupted routine data flows; at the same time, displacement has remained large and mobile, shifting catchments and denominators. Under such conditions, the value of simulation lies less in precise point prediction and more in converting noisy inputs into threshold-based options that map to incident objectives and orders. Documented counts of attacks and current displacement totals indicate this is not a transient problem. Hence, making simulation a routine part of PHEOC planning is a practical response rather than a research ideal.

Interpreting our results from data sources, we see that the literature supports a pragmatic triangulation strategy. Social media advertising data, mobile device mobility traces, and satellite-based car counts introduce bias. However, when combined, they move faster than registries and can track subnational population change closely enough to steer outreach and staging. The implication for practice is not to replace administrative data but to treat these feeds as early signals that are folded into the PHEOC information function with explicit caveats and periodic checks against surveys or official tallies.

Our results also indicated that disease risks around Ukraine pull in different directions: a surge of measles across the WHO European Region in 2024 demands catch-up vaccination and outreach, while the formal closure of Ukraine's 2021 cVDPV2 polio outbreak shows that targeted campaigns can succeed even during conflict. Rather than contradicting each other, these signals point to modular use of simulation: short-term demand forecasts and micro-planning support for measles; monitoring and contingency scenarios to maintain polio immunity in mobile populations.

For tuberculosis, our results suggest shifting the modelling question from «how many infections?» to «where does the care cascade break under strain?». Regional surveillance highlights the continued burden of MDR/RR-TB and the risk posed by treatment interruptions [35]. In that context, scenario comparisons that test adherence support, drug supply continuity, and cross-border referral are more decision-relevant than transmission-only outputs.

Environmental disruption after the Kakhovka Dam breach sharpens this orientation toward scenarios. Studies document altered hydrology and likely mobilization of contaminants, but the operational issue is whether current WASH measures and surveillance remain adequate under different access constraints. The results, therefore, support embedding WASH-sensitive scenarios, safe-water distribution, vector control, and intensified testing into the planning cycle for the affected oblasts.

The enablers and barriers we identified also interpret the results in practical terms. Scoping reviews of PHEOCs consistently flag data management, staffing, and information-sharing as limiting factors; our mapping to ICS artefacts directly responds to those gaps, making it easier for numbers to become orders. Likewise, the prominence of simulation exercises and after-action reviews in current ECDC and WHO guidance matches the result that integration improves when teams rehearse thresholds and hand-offs before activation.

The results on communication argue for disciplined uncertainty practices. WHO's 2024–2025 guidance clearly states that confidence grows when leaders and the public hear what is known, what is unknown, and what would change the advice. Translating that into one-page briefs for the Planning meeting, paired with a decision log, turns a general principle into a repeatable habit that supports accountability and rapid revision.

Conclusions

This chapter shows that epidemic simulation has the greatest impact when treated as a routine input to command work rather than an occasional technical add-on. This framing is feasible and necessary in the Ukraine context, marked by attacks on health care, sustained displacement, and environmental shocks: near-term ensembles support operations. At the same time, scenario exercises guide strategy under shifting constraints.

The paper consolidates recent literature into a coherent, operational integration template for crisis programs. It contributes a dual-track approach that pairs short-horizon ensembles with medium-term scenario hubs; a displacement-aware data strategy that triangulates social media, mobility, and satellite signals; an explicit equity checkpoint embedded in ICS/PHEOC cycles; and a direct mapping of model outputs to incident artefacts.

The evidence-informed checklist can be adopted immediately: designate a modelling liaison in Planning, issue weekly one-page ensemble briefs, run scheduled scenario cycles linked to triggers, maintain a decision log, and stand up minimum data feeds for surveillance, capacity, displacement, and WASH. These steps reduce the friction between analysis and orders, improve accountability, and are transferable to other high-constraint settings. Future work should evaluate how these practices change decisions and outcomes and standardise reporting for operational modelling.

This study was funded by the National Research Foundation of Ukraine in the framework of the research project 2023.03/0197 on the topic «Multidisciplinary study of the impact of emergency situations on the infectious diseases spreading to support management decision making in the field of population biosafety».

References

1. IHR, «Public health emergency operations centres», *World Health Organization – Regional Office for the Eastern Mediterranean*, 2025. <https://www.emro.who.int/international-health-regulations/areas-of-work/public-health-emergency-operations-center.html> (accessed Aug. 10, 2025).
2. The Lancet Commission on Strengthening the Use of Epidemiological Modelling of Emerging and Pandemic Infectious Diseases, «How modelling can better support public health policy making: the Lancet Commission on Strengthening the Use of Epidemiological Modelling of Emerging and Pandemic Infectious Diseases», *The Lancet*, vol. 403, no. 10429, pp. 789–791, Dec. 2023, doi: [https://doi.org/10.1016/s0140-6736\(23\)02758-7](https://doi.org/10.1016/s0140-6736(23)02758-7)

3. E. Howerton *et al.*, «Evaluation of the US COVID-19 Scenario Modeling Hub for informing pandemic response under uncertainty», *Nature Communications*, vol. 14, no. 1, p. 7260, Nov. 2023, doi: <https://doi.org/10.1038/s41467-023-42680-x>
4. M. Jit and A. R. Cook, «Informing Public Health Policies with Models for Disease Burden, Impact Evaluation, and Economic Evaluation», *Annual Review of Public Health*, vol. 45, no. 1, pp. 133–150, Oct. 2023, doi: <https://doi.org/10.1146/annurev-publhealth-060222-025149>
5. T. Allen and R. Spencer, «Barriers and Enablers to Using an Emergency Operations Center in Public Health Emergency Management: A Scoping Review», *Disaster Medicine and Public Health Preparedness*, vol. 17, p. e407, Jan. 2023, doi: <https://doi.org/10.1017/dmp.2023.50>
6. WHO, «Three years of war: rising demand for mental health support, trauma care and rehabilitation», *Who.int*, Feb. 24, 2025. <https://www.who.int/europe/news/item/24-02-2025-three-years-of-war-rising-demand-for-mental-health-support-trauma-care-and-rehabilitation> (accessed Aug. 10, 2025).
7. UNHCR, «Ukraine emergency», *UNHCR*, Jun. 2023. <https://www.unhcr.org/emergencies/ukraine-emergency> (accessed Aug. 10, 2025).
8. World Health Organization, «European Region reports highest number of measles cases in more than 25 years – UNICEF, WHO/Europe», *Who.int*, Mar. 13, 2025. <https://www.who.int/europe/news/item/13-03-2025-european-region-reports-highest-number-of-measles-cases-in-more-than-25-years---unicef--who-europe> (accessed Aug. 10, 2025).
9. World Health Organization, «Polio outbreak in Ukraine closed – a success story for public health despite extreme challenges of war», *www.who.int*, 2023. <https://www.who.int/europe/news/item/21-09-2023-polio-outbreak-in-ukraine-closed-a-success-story-for-public-health-despite-extreme-challenges-of-war> (accessed Aug. 10, 2025).
10. World Health organization, «Global Tuberculosis Report 2024», *Who.int*, 2024. <https://www.who.int/publications/i/item/9789240101531> (accessed Aug. 10, 2025).
11. O. Shumilova *et al.*, «Environmental effects of the Kakhovka Dam destruction by warfare in Ukraine», *Science*, vol. 387, no. 6739, pp. 1181–1186, Mar. 2025, doi: <https://doi.org/10.1126/science.adn8655>
12. D. R. Leasure *et al.*, «Nowcasting Daily Population Displacement in Ukraine through Social Media Advertising Data», *Population and Development Review*, vol. 49, no. 2, pp. 231–254, Apr. 2023, doi: <https://doi.org/10.1111/padr.12558>
13. Y. Shibuya, N. Jones, and Y. Sekimoto, «Assessing internal displacement patterns in Ukraine during the beginning of the Russian invasion in 2022», *Scientific Reports*, vol. 14, no. 1, p. 11123, May 2024, doi: <https://doi.org/10.1038/s41598-024-59814-w>
14. World Health Organization, «Public Health Emergency Operations Centre Network (EOC-NET)», *www.who.int*, 2021. <https://www.who.int/groups/eoc-net> (accessed Aug. 10, 2025).
15. CDC, «Emergency Operations Centers and Incident Management Structure», *Field Epidemiology Manual*, Sep. 05, 2024. <https://www.cdc.gov/field-epi-manual/php/chapters/eoc-incident-management.html> (accessed Aug. 10, 2025).
16. K. Sherratt *et al.*, «Predictive performance of multi-model ensemble forecasts of COVID-19 across European nations», *eLife*, vol. 12, p. e81916, Apr. 2023, doi: <https://doi.org/10.7554/elife.81916>
17. S. Jung *et al.*, «Potential impact of annual vaccination with reformulated COVID-19 vaccines: Lessons from the US COVID-19 scenario modeling hub», *PLoS Medicine*, vol. 21, no. 4, p. e1004387, Apr. 2024, doi: <https://doi.org/10.1371/journal.pmed.1004387>

18. M. C. Runge *et al.*, «Scenario design for infectious disease projections: Integrating concepts from decision analysis and experimental design», *Epidemics*, vol. 47, p. 100775, May 2024, doi: <https://doi.org/10.1016/j.epidem.2024.100775>
19. C. J. Owek *et al.*, «Lessons learned from COVID-19 modelling efforts for policy decision-making in lower- and middle-income countries», *BMJ Global Health*, vol. 9, no. 11, p. e015247, Nov. 2024, doi: <https://doi.org/10.1136/bmjgh-2024-015247>
20. K. R. Moran, T. Lopez, and S. Y. Del Valle, «The future of pandemic modeling in support of decision making: lessons learned from COVID-19», *BMC Global and Public Health*, vol. 3, no. 1, p. 24, Mar. 2025, doi: <https://doi.org/10.1186/s44263-025-00143-z>
21. M.-C. Rufener, F. Ofli, M. Fatehikia, and I. Weber, «Estimation of internal displacement in Ukraine from satellite-based car detections», *Scientific Reports*, vol. 14, no. 1, p. 31638, Dec. 2024, doi: <https://doi.org/10.1038/s41598-024-80035-8>
22. N. Wycoff, L. O. Singh, A. Arab, K. M. Donato, and Helge Marahrens, «The digital trail of Ukraine's 2022 refugee exodus», *Journal of Computational Social Science*, vol. 7, pp. 2147–2193, Jul. 2024, doi: <https://doi.org/10.1007/s42001-024-00304-4>
23. European Centre for Disease Prevention and Control, «Measles on the Rise Again in Europe: Time to Check Your Vaccination Status», *European Centre for Disease Prevention and Control*, Mar. 11, 2025. <https://www.ecdc.europa.eu/en/news-events/measles-rise-again-europe-time-check-your-vaccination-status> (accessed Aug. 10, 2025).
24. European Centre for Disease Prevention and Control, «Tuberculosis surveillance and monitoring in Europe 2024 – 2022 data», *www.ecdc.europa.eu*, Mar. 21, 2024. <https://www.ecdc.europa.eu/en/publications-data/tuberculosis-surveillance-and-monitoring-europe-2024-2022-data> (accessed Aug. 10, 2025).
25. D. Jiang *et al.*, «The biogeochemical response of the north-western Black Sea to the Kakhovka Dam breach», *Communications Earth & Environment*, vol. 6, no. 1, p. 185, Mar. 2025, doi: <https://doi.org/10.1038/s43247-025-02153-z>
26. United Nations, «Rapid Environmental Assessment of Kakhovka Dam Breach», 2023. Accessed: Aug. 10, 2025. [Online]. Available: https://unepdhi.org/wp-content/uploads/sites/2/2023/11/UNEP_Kakhovka_Dam_Breach_Ukraine_Assessment.pdf
27. European Centre for Disease Prevention and Control, «Strengthening health resilience through simulation exercises and after-action reviews», *European Centre for Disease Prevention and Control*, Jun. 13, 2025. <https://www.ecdc.europa.eu/en/news-events/strengthening-health-resilience-through-simulation-exercises-and-after-action-reviews> (accessed Aug. 19, 2025).
28. A. Chambers, I. Khan, R. Repchuck, S. Muir, H. Hanson, and Y. Khan, «Enhancing the design, conduct and evaluation of public health emergency preparedness exercises: a rapid review», *BMC Public Health*, vol. 25, no. 1, p. 2366, Jul. 2025, doi: <https://doi.org/10.1186/s12889-025-23270-6>
29. R. Moyal-Smith, D. J. Barnett, E. Toner, J. A. Marsteller, and C. T. Yuan, «Embedding Equity into the Hospital Incident Command System: A Narrative Review», *The Joint Commission Journal on Quality and Patient Safety*, vol. 50, no. 1, pp. 49–58, Oct. 2023, doi: <https://doi.org/10.1016/j.jcjq.2023.10.011>
30. R. Moyal-Smith, J. A. Marsteller, D. J. Barnett, P. Kent, T. Purnell, and C. T. Yuan, «Centering Health Equity in the Implementation of the Hospital Incident Command System: A Qualitative Case Comparison Study», *Disaster medicine and public health preparedness*, vol. 18, pp. 1–24, Feb. 2024, doi: <https://doi.org/10.1017/dmp.2024.20>

31. World Health Organization, «Communicating uncertainty in health emergencies – guidance and tips», *Who.int*, 2024. <https://www.who.int/europe/publications/m/item/communicating-uncertainty-in-health-emergencies-guidance-and-tips> (accessed Aug. 10, 2025).
32. L. Hadley, C. Rich, A. Tasker, O. Restif, and S. Funk, «How does policy modelling work in practice? A global analysis on the use of epidemiological modelling in health crises», *PLOS Global Public Health*, vol. 5, no. 6, p. e0004675, Jun. 2025, doi: <https://doi.org/10.1371/journal.pgph.0004675>
33. E. Y. Cramer *et al.*, «Evaluation of individual and ensemble probabilistic forecasts of COVID-19 mortality in the United States», *Proceedings of the National Academy of Sciences*, vol. 119, no. 15, p. e2113561119, Apr. 2022, doi: <https://doi.org/10.1073/pnas.2113561119>
34. COVID-19 Forecast Hub, «Ensemble model – COVID 19 forecast hub», *covid19forecasthub.org*, 2021. <https://covid19forecasthub.org/doc/ensemble/> (accessed Aug. 10, 2025).
35. M. Sartipi, A. Aminafshar, A. Pakzad, M. Shafiei, H. Farahmandnia, and A. Tavan, «Evaluating the efficiency of emergency operation centers during pandemics: a cross-sectional study with operational managers in Iran», *Frontiers in Public Health*, vol. 13, p. 1511932, May 2025, doi: <https://doi.org/10.3389/fpubh.2025.1511932>

INFRASTRUCTURE PROJECTS: ASSESSMENT OF ANTHROPOGENIC IMPACT ON RIVER ECOSYSTEMS

Danshyna S., Podorozhko K.

The aim of the study is to improve the accuracy of assessing anthropogenic impacts on rivers by developing an integrated approach that combines functional modeling and satellite monitoring. The methodological foundations for assessing anthropogenic pressure on river ecosystems within infrastructure projects are considered. Using the IDEF0 methodology, a functional model of the assessment process is proposed, which structures the sequence of actions, input and output data for monitoring and analyzing anthropogenic impacts on rivers. Particular attention is paid to the processing of data on anthropogenic factors, in particular hydrological, morphological, and Earth remote sensing data.

Introduction

Infrastructure projects (e.g., construction of water supply, sewage, transport, and energy facilities, residential and social buildings, etc.) which implement strategic efforts to achieve social, economic, and environmental sustainability goals aimed at solving long-term problems, become a means of creating and realizing values with long-term consequences for the environment and society. Such projects are expensive, controversial, and complex to manage, requiring mandatory assessment of their impact on natural resources, ecosystem services, and social services to prevent overspending, social conflicts, and reputational losses, ensuring project sustainability [1, 2]. Therefore, a comprehensive approach to assessing the environmental impact of infrastructure projects is needed, combining traditional assessments of cost, time, and quality [1].

Intensive infrastructure development – transport, energy, water supply, housing construction – is increasingly affecting the state of water ecosystems, particularly rivers, which are sensitive to anthropogenic changes. The construction of facilities within river basins is accompanied by [3]:

1. Changes in riverbed processes, as construction affects water flow velocity, sediment transport, and riverbed erosion, which in turn leads to a decrease in the regenerative capacity of riverbeds.

2. Pollution of rivers, as the discharge of industrial waste, agricultural chemicals, and urban wastewater significantly deteriorates water quality. Water pollution with heavy metals, pesticides, and nitrates leads to the degradation of ecosystems and the deterioration of living conditions for many aquatic species.

3. Changes in the hydrological regime, in particular due to agricultural land reclamation, the use of water resources for industrial needs, and the straightening of river channels.

In this regard, there is a growing need for reliable tools that allow for a scientifically sound assessment of the impact of infrastructure projects on the aquatic environment at the planning and decision-making stages [4].

Functional model of the assessment process for anthropogenic impact on rivers

When developing infrastructure projects, it has become common practice to consider the various consequences of decisions, which are characterized not only by direct profits but also by non-monetary aspects. In particular, environmental, health, and safety issues are receiving increasing attention to prevent both short-term and long-term damage [5]. The anthropogenic impact on rivers is a multifaceted problem, the theoretical, methodological, and methodological aspects of which are addressed in scientific works by geographers, hydrochemists, hydromeliorators, and ecologists Vyshnevsky P.F., Kirilyuk O.V., Levkivsky S.S., Likho O.A., Myskovets I.Ya., Morokova V.V., Rybalova O.V., Solovey T.V., Tymchenko Z.V., Khilchevsky V.K., Tsvetova O.V., Yasenchuk T.O., Yatsyk A.V., and others. Various methods have been developed that can be used to identify weaknesses in infrastructure projects that require technical measures or to compare different options for prioritizing resource investments. However, their scope of application varies between methods. For example, infrastructure project risk assessment focuses on random problems that may arise in river systems, while life cycle assessment (LCA) identifies long-term and sustainable impacts on river basins. Ideally, several assessment methods should be used to take into account the impact on different areas when making decisions [1, 5]. However, in practice, project managers face difficulties:

- when collecting the necessary data for the implementation of non-traditional approaches to project assessment, such as LCA;
- when combining assessment results obtained using different methods during decision-making;
- when systematizing tools, data, and measures at different stages of the project life cycle.

Therefore, scientists increasingly turn to Earth remote sensing (ERS) methods in the study and assessment of anthropogenic impacts of infrastructure projects. For example, the use of satellite data and aerial photography allows assessing the

scale and dynamics of changes in river systems under the influence of human activity. In scientific publications on this topic, ERS data are used for [2–4]:

1. Monitoring changes in river channels using multi-year series of satellite images, which allows identifying changes within the channel, displacement of channel forms, and changes in coastlines.

2. Assessing the impact of the construction of hydraulic structures (dams, weirs, canals), when ERS data make it possible to clearly track the effects of sediment retention upstream, bank erosion in the lower reaches of rivers, and artificial changes in the hydrodynamics of river systems, as well as to identify the consequences of flooding or erosion of banks, etc.

3. Determination of changes in river systems and degradation of natural river landscapes due to agricultural development, changes in drainage, and/or land reclamation. Here, satellite images are used to track changes in the area of irrigated land, the expansion of agricultural land, the drainage of coastal areas, and the increase in water abstraction from rivers.

4. Investigation of erosion processes and sedimentation over large areas, where high-resolution images are used to analyze coastal erosion, identify areas most susceptible to sedimentation processes, track changes in channel dynamics, etc.

The combination of modern geographic information systems (GIS) and ERS data creates new approaches for assessing the impact of infrastructure projects on the aquatic environment. However, in practice, there is often no comprehensive solution that would allow not only to record changes, but also to proactively mitigate the effects of impact, taking into account spatial, temporal, and regulatory parameters. That is why, at the planning, assessment, and implementation stages of an infrastructure project, it is important to have a tool that allows [4]:

- model scenarios of impact on the aquatic environment,
- assess the level of risk from human activities,
- visualize the spatial location of critical areas (pollution, riverbed degradation, etc.),
- make management decisions based on data rather than intuition or «average estimates».

In recent years, such decision support systems have become increasingly important. This trend will continue in the future, but most likely on a much larger scale. Although until recently the focus was on individual software tools and services, their integration into engineering workflows is becoming a new and complex area of research and development [6], requiring a systematic approach that takes into account the interrelationships between natural, technical, social, and managerial components.

The key point here is to accurately define the necessary data and information, tools, and modeling mechanisms at different stages of evaluation and to present the entire process in the form of a business model.

An effective tool for developing a decision support system is the IDEF0 functional modeling methodology. It presents the system as a business model of comprehensive information analysis [5–7], which is commonly used in several infrastructure project evaluation methods, with a particular focus on the environmental assessment of its life cycle and the assessment of risks during project implementation (Fig. 1).

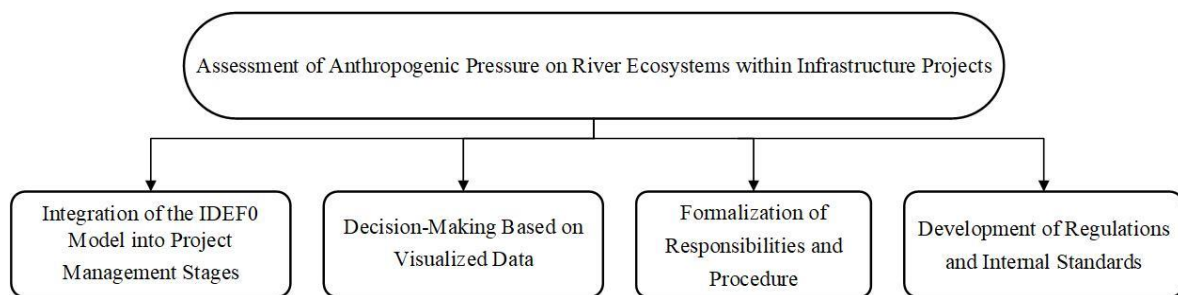


Fig. 1. Stages of assessment of the anthropogenic impact on river ecosystems within infrastructure projects

Thus, in the functional model, environmental, health, and safety aspects are considered as new project evaluation criteria alongside traditional economic and technical indicators. This approach allows linking the logic of environmental monitoring processes with project implementation processes, visualizing information and functional links between decision-making stages, and identifying points of influence where environmental risks can be minimized [7].

Fig. 2 shows the IDEF0 model of the process of assessing the anthropogenic impact of infrastructure projects on rivers. In accordance with the provisions of the functional modeling methodology, this model defines the structure of information technology for assessing the anthropogenic impact on the state of rivers. It provides a systematic and transparent description of the project assessment stages, when the manager must take into account various types of requirements for the conservation of water and biological resources and the maintenance of the ecological status of river ecosystems, the availability of relevant design constraints and resources at different stages of design, etc.

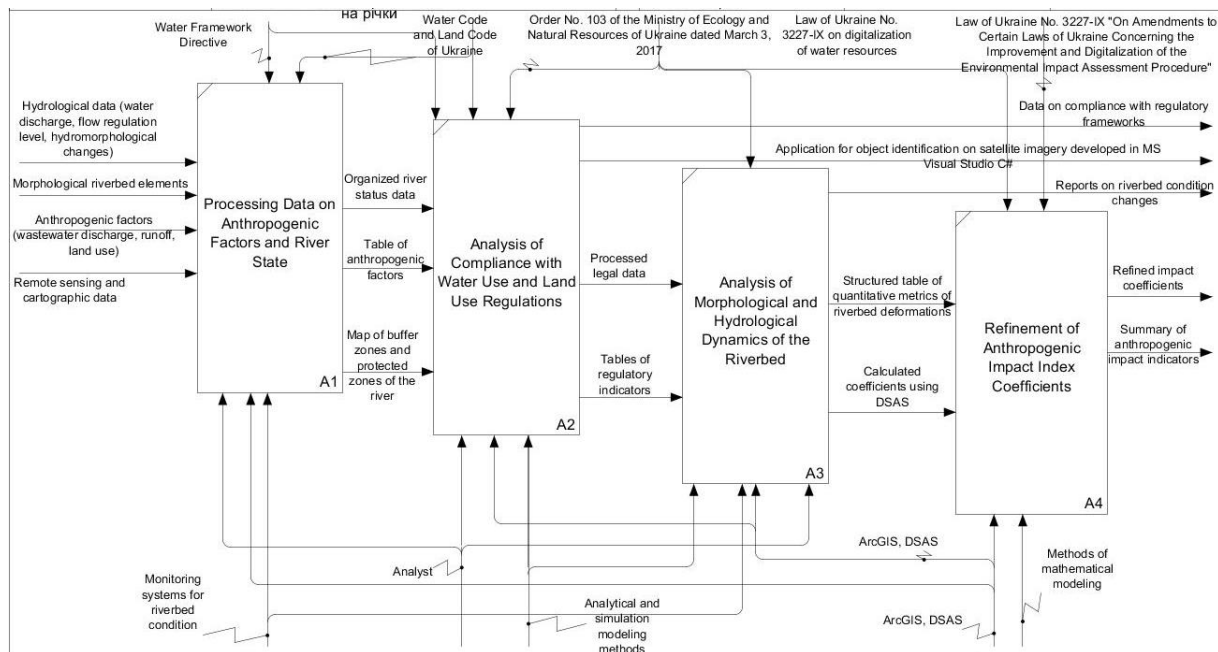


Fig. 2. IDEF0 model as a structure of information technology for assessment of anthropogenic impact on river ecosystems within infrastructure projects

The proposed IDEF0 model visualizes the procedure for obtaining assessments of anthropogenic impact on rivers, taking into account various aspects of data acquisition at the stages of the infrastructure project life cycle (Table 1). It provides an understanding of the characteristics and volumes of data required for the formation of assessments [6, 8] and enables the creation of regulations and internal project standards. For example, based on the IDEF0 model, it is possible to develop:

- regulations for environmental support of the project (when and which process blocks are involved at the stage of assessment formation);
- checklists for project managers (what data is needed to obtain assessments, when they should be analyzed);
- algorithms for decision-making and developing alternatives (for example, in cases of exceeding the impact level).

Such a system contributes to increasing the transparency of decisions, timely identification of environmental risks, and optimization of control procedures. Thanks to its integration into the planning, implementation, and monitoring stages, the IDEF0 model helps reduce uncertainty in management decision-making, which in turn improves project management efficiency and reduces the anthropogenic load on river ecosystems [7, 8].

Table 1

Using the IDEF0 model in project management stages

Project life cycle stage	Use of the IDEF0 model
Initiation	– provision of basic environmental data on the river basin area; – risk analysis at the site selection stage.
Planning	– modeling of the impact of alternative project solutions; – development of environmental protocols.
Implementation	– assessment of the actual impact on the river ecosystem; – monitoring of project deviations; – development of corrective measures.
Completion	– generation of a final project report; – assessment of the residual load on river ecosystems; – recommendations for the future.

In the changing conditions of infrastructure project implementation, uncertainty in the distribution of functions among project participants, the absence of clearly defined procedures for action in critical environmental situations, and insufficient integration between technical and environmental solutions can lead to irreversible damage to water ecosystems. Formalization based on the proposed IDEF0 model allows establishing transparent links between functional management blocks and specific performers, regulatory constraints, and documentation. This ensures not only the effectiveness of management decisions but also their compliance with environmental legislation and international standards [8, 9]. In conditions of environmental instability and growing demands for sustainable development, effective management of environmental projects requires a clear understanding of who makes decisions, when, and on the basis of what data [9, 10]. Each block of the IDEF0 model can be linked to specific departments of organizations or institutions (State Water Agency, design institutes, environmental departments, etc.), documents (technical specifications, water body passport, etc.), regulatory restrictions (water protection zones, maximum permissible concentrations, etc.).

The IDEF0 model also provides elements for visualizing the obtained assessments using GIS tools. Traditional approaches based solely on text documents are insufficient in complex interdisciplinary tasks where spatial, temporal, and risk-oriented indicators must be taken into account [8–10]. Visualization of data through GIS maps, dynamic forecasting models, or analytical dashboards makes it possible to increase the transparency and soundness of decisions, simplify communication between participants, and minimize the influence of the human factor.

Thus, the proposed IDEF0 model can be viewed as an «interaction architecture» between data, processes, and decisions. Its use allows environmental issues to be integrated into the actual process of infrastructure project management –

from analytics to practical environmental risk management in the context of anthropogenic pressure on rivers.

Processing data on anthropogenic impact factors on the state of river ecosystems

According to Fig. 2, at the first stage of assessing anthropogenic impact, data of various types relating to the state of rivers should be processed. This stage is based on the integration of information from various information resources, among which hydrological, morphological indicators, and cartographic data play a key role.

Hydrological indicators include [11]: water consumption, the level of flow regulation, and changes in the hydromorphological state of the riverbed. Morphological data characterize the configuration and structural parameters of riverbed forms and allow the identification of areas of intense erosion or sediment accumulation [12, 13]. Anthropogenic factors of human activity, such as the degree of coastal land cultivation, urbanization, and military operations, are considered separately [14].

An important source of information is Earth remote sensing (ERS) data obtained from high-resolution satellite images and aerial photography. The use of multispectral images, in particular those obtained from Landsat satellites (Table 2) and radar images, makes it possible to detect changes in the water surface, determine the degree of siltation and fragmentation of aquatic ecosystems, and perform automated classification of objects using GIS, in particular

ArcGIS and specialized processing software (in particular, using algorithms in MS Visual Studio C#) [11, 12, 14].

At the processing stage, the input ERS data undergoes preliminary cleaning, georeferencing, and integration into a single information and analytical database [10, 11].

The results obtained are used to construct buffer zones for hydrological stations, form tables of hydrological observations, and create maps of the spatial distribution of anthropogenic loads. This enables further analysis in the subsequent stages of assessment – verification of compliance with regulatory requirements, modeling of riverbed morphodynamics, and refinement of anthropogenic impact indices [14].

Table 2

Information about Landsat 8, 7, 5 satellite channels [10, 11]

Description of channel ranges	Landsat 8		Landsat 7		Landsat 5	
	Channel number	Wavelength, μm	Channel number	Wavelength, μm	Channel number	Wavelength, μm
Ultraviolet (UV)	1	0,433 – 0,453	–	–	–	–
Blue	2	0,450 – 0,515	1	0,450 – 0,515	1	0,450 – 0,515
Green	3	0,525 – 0,600	2	0,525 – 0,605	2	0,525 – 0,605
Red	4	0,630 – 0,680	3	0,630 – 0,690	3	0,630 – 0,690
Near IR (NIR)	5	0,845 – 0,885	4	0,750 – 0,900	4	0,750 – 0,900
Shortwave IR 2 (SWIR 2)	6	1,560 – 1,660	5	1,550 – 1,750	5	1,550 – 1,750
Shortwave IR 3 (SWIR 3)	7	2,100 – 2,300	7	2,090 – 2,350	7	2,090 – 2,350
Panchromatic (PAN)	8	0,500 – 0,680	8	0,520 – 0,900	–	–
Shortwave IR (SWIR)	9	1,360 – 1,390	–	–	–	–
Far IR (FIR 1)	10	10,30 – 11,30	6-1	10,30 – 11,30	6	10,40 – 12,5
Far IR 2 (FIR 2)	11	11,50 – 12,50	6-2	10,30 – 11,30	–	–

Thus, the use of spatial analysis methods and ERS data at the early stages of an infrastructure project's life cycle ensures a more informed assessment of the environmental status of river systems, which in turn improves planning quality, minimizes risks, and optimizes management decisions. This approach contributes to more effective project implementation and enables the achievement of sustainable development goals in the field of water ecosystems.

Conclusions

To improve the accuracy of assessing anthropogenic impacts on river ecosystems within infrastructure projects, a functional model of the process has been developed that combines system modeling using the IDEF0 methodology and modern Earth remote sensing data processing tools. The functional model clearly structures the assessment process, defining the relationships between input data, analysis tools, and expected results.

The proposed approach involves the integration of various types of data, including hydrological, morphological, and cartographic data with multispectral satellite images, ensuring improved spatial and temporal accuracy and relevance of assessments. The use of GIS and specialized software allows automating the

detection of changes in riverbed conditions, assessing the degree of anthropogenic load, and timely identifying critical areas.

The practical significance of the results obtained lies in the possibility of applying the developed model to improve water resource management, plan environmental protection measures, and minimize the negative impacts of infrastructure projects on river ecosystems. There is potential for scaling up to other basin systems and integration into national and international environmental monitoring systems. Issues of environmental support for infrastructure projects, the development of regulatory and methodological documents, and the formation of effective environmental policy in the water sector are being addressed.

The work was carried out with the support of the Ministry of Education and Science of Ukraine (state registration numbers of projects 0125U000573, 0125U001680).

Reference

1. Danshyna S., Podorozhko K. Spatial assessment of infrastructure projects. *Decision Support Systems in Project and Program Management* : collective monograph / ed. by I. Linde. Riga : European University Press, ISMA University of Applied Sciences, 2024. P. 41 – 52. DOI: 10.30837/MMP.2024.041
2. Impacts of Infrastructure Developments on Ecosystem Services / E. Y. Menteşe, T. Kaya, A. Erol, M. F. Türker. *Frontiers in Environmental Science*. 2021. Vol. 9. Article 614752. DOI: 10.3389/fenvs.2021.614752
3. Assessing anthropogenic impacts on river ecosystems using PLS regression and GIS: A case study in the Ave River basin (Portugal) / A. R. L. Ferreira, J. P. Nunes, N. Abrantes, F. Gonçalves. *Science of the Total Environment*. 2017. Vol. 579. P. 1094 – 1104. DOI: 10.1016/j.scitotenv.2016.11.144
4. Rousseau A. N., Mailhot A., Turcotte R. Integration of hydrological modelling and GIS for watershed management. *Journal of Hydrologic Engineering*. 2000. Vol. 5, issue 2. P. 174 – 182. DOI: 10.1061/(ASCE)1084-0699(2000)5:2(174)
5. Analysis and Modeling of Information Required for Process Assessment on Environment, Health, and Safety by IDEF0 and UML / Y. Kikuchi et al. *Computer Aided Chemical Engineering*. 2012. Vol. 31. P. 1392 – 1396. DOI: 10.1016/B978-0-444-59506-5.50109-7
6. Schneider R. A., Marquardt W. Information technology support in the chemical process design life cycle. *Chemical Engineering Science*. 2002. Vol. 57, issue 10. P. 1763 – 1792. DOI: 10.1016/S0009-2509(02)00075-1
7. IDEF0 activity modeling for integrated process design considering environmental, health and safety (EHS) aspects / M. Hirao, H. Sugiyama, U. Fischer, K. Hungerbühler. *Computer Aided Chemical Engineering*. 2008. Vol. 25. P. 1065 – 1070. DOI: 10.1016/S1570-7946(08)80184-8

8. Too E. G., Weaver P. The management of project management: A conceptual framework for project governance. *International Journal of Project Management*. 2014. Vol. 32, issue 8. P. 1382 – 1394. DOI: 10.1016/j.ijproman.2013.07.006
9. Project Management Institute. *A Guide to the Project Management Body of Knowledge (PMBOK® Guide)*. Project Management Institute, Inc., 2021. 478 p.
10. Kumar G., Bali R., Agarwal A. K. GIS integration of remote sensing and electrical data for hydrological exploration: a case study of Bhakar watershed, India. *Hydrological Sciences Journal*. 2009. Vol. 54, issue 5. P. 949 – 960. DOI: 10.1623/hysj.54.5.949
11. Remote Sensing Satellite. *Sentinel-2 Series*. Available at: <https://ecoruspace.me/Sentinel+2.html> (accessed 22 March 2025).
12. Unraveling the Impacts of River Network Connectivity on Ecological Quality Dynamics at a Basin Scale / X. Li et al. *Remote Sensing*. 2024. Vol. 16, issue 13. Article 2370. DOI: 10.3390/rs16132370
13. Herzog A., Hellwig J., Stahl K. An investigation of anthropogenic influences on hydrologic connectivity using model stress tests. *Hydrology and Earth System Sciences*. 2024. Vol. 28, P. 4065 – 4083. DOI: 10.5194/hess-28-4065-2024
14. Danshina S. Yu., Podorozhko K. D. Method of spatial analysis of riverbed deformation based on combined data. *Restoration of ecosystems damaged by military operations: Ukrainian and European challenges: abstracts of reports from the International Scientific and Practical Conference REDMO-2024*, Kyiv, November 6, 2024. K., 2024. P. 20–25.

MULTI-CRITERIA DECISION-MAKING APPROACHES FOR IT OUTSOURCING PROJECT PORTFOLIO MANAGEMENT

Dobrytskyi D., Fonarova T.

The article provides an in-depth review of key multi-criteria decision-making (MCDM) approaches relevant to IT outsourcing project portfolio management. The discussion covers the classification of methods – from hierarchical and distance-based techniques to outranking, utility/efficiency, and hybrid approaches – with a focus on their strengths, limitations, mathematical requirements, and applicability in IT outsourcing environments. Comparative tables summarise advantages and disadvantages, typical application scenarios, and decision-support criteria. The paper also examines methodological aspects of selecting an MCDM method, considering data types, portfolio scale, uncertainty, resource constraints, and integration with collaborative decision-making tools. Additional frameworks for method selection are proposed, aiming to support IT outsourcing companies in aligning portfolio decisions with strategic goals.

Introduction

Over the last decade, the IT outsourcing industry has shown stable growth, accompanied by an intensification of competition. Companies aiming to maintain or strengthen their market positions can no longer rely solely on financial indicators when evaluating potential projects. Experience demonstrates that focusing exclusively on revenue or profitability leads to the gradual erosion of key competencies and increased portfolio vulnerability to market fluctuations. This challenge is particularly relevant in emerging IT service economies, where export revenue from the sector constitutes a significant portion of national GDP. In such contexts, competitive pressures demand that companies select projects based not only on short-term profit but also on strategic alignment, capability development, risk control, and optimal resource allocation.

Multi-Criteria Decision-Making (MCDM) methods provide a structured framework to address this challenge. Unlike intuitive judgment or single-criterion models, MCDM enables decision-makers to evaluate multiple heterogeneous factors simultaneously, reduce subjectivity, and ensure that choices remain aligned with long-term corporate strategy.

The selection of a specific method depends on various factors, including the nature of available data, the complexity of the portfolio, the number of stakeholders involved, and the resources available for model development. In the IT outsourcing context, methods must be flexible enough to handle both quantitative and qualitative criteria and support collaborative decision-making processes.

1. Classification of MCDM Methods

In project portfolio management research, MCDM methods are commonly grouped into categories that reflect their conceptual foundation and computational logic. This section expands on each method group, outlining its main principles, mathematical basis, and practical implications for IT outsourcing companies.

Table 1

Classification of MCDM methods

Method group	Examples	Core principle	Data type	Typical IT outsourcing tasks
Hierarchical	AHP, ANP	Pairwise comparison and criteria hierarchy	Quantitative , qualitative	Determining criteria weights
Distance-based	TOPSIS, VIKOR	Proximity to ideal solution	Quantitative	Project ranking
Outranking	ELECTRE, PROMETHEE	Establishing relative preferences	Quantitative , qualitative	Vendor/partner evaluation
Utility/ Efficiency-based	MAUT, DEA	Integrated utility or efficiency assessment	Quantitative	Resource efficiency evaluation
Hybrid	AHP–TOPSIS, ANP–VIKOR	Combination of methods	Mixed	Comprehensive project selection

1.1 Hierarchical methods (AHP, ANP) – decompose complex decision problems into levels: a goal at the top, criteria in the middle, and alternatives at the bottom [6]. Decision-makers perform pairwise comparisons of elements within the same level using a numerical scale (e.g., Saaty’s 1–9 scale) to express relative importance. The comparison matrices are processed to derive priority vectors (weights), and a consistency ratio (CR) is calculated to verify logical coherence. In IT outsourcing, these methods are useful for criteria weight determination when selecting contracts or allocating resources.

1.2 Distance-based methods (TOPSIS, VIKOR) – rank alternatives by calculating their relative closeness to an ideal solution and their distance from an anti-ideal solution [4]. For example, in TOPSIS, each criterion is normalized, weighted, and then used to compute Euclidean distances to the ideal and anti-ideal points. The final score is the closeness coefficient.

1.3 Outranking methods (ELECTRE, PROMETHEE) – compare alternatives pairwise to determine whether one outranks another based on concordance and discordance indices [5, 7]. In IT outsourcing, these methods are valuable when qualitative assessments are as important as quantitative measures.

1.4 Utility/Efficiency-based methods (MAUT, DEA) – transform performance on each criterion into a utility score or measure the efficiency of each alternative as the ratio of weighted outputs to weighted inputs [8, 9].

1.5 Hybrid methods (AHP–TOPSIS, ANP–VIKOR, Fuzzy MCDM) – combine complementary strengths of different approaches [3]. For IT outsourcing, hybrid methods are attractive when both strategic alignment and numeric performance indicators must be considered in one framework.

2. Comparative Analysis for IT Outsourcing

The diversity of multi-criteria decision-making methods available to managers in the IT outsourcing sector reflects the broad spectrum of decision contexts encountered in practice. While all MCDM approaches share the common goal of integrating multiple evaluation criteria into a single, coherent assessment, they differ significantly in their mathematical foundations, data requirements, interpretability, and suitability for specific operational environments. Understanding these differences is critical for ensuring that the selected method not only yields a robust ranking of alternatives but also fits the strategic and operational realities of a given company.

In the context of IT outsourcing project portfolio management, hierarchical methods such as the Analytic Hierarchy Process (AHP) and the Analytic Network Process (ANP) are often appreciated for their conceptual clarity and ability to capture subjective judgments in a structured way [6]. They enable decision-makers to decompose complex problems into manageable components, assigning weights to criteria through systematic pairwise comparisons. This transparency makes them particularly effective for strategic discussions in which consensus among stakeholders is important. However, their reliance on human judgment can become a limitation when the number of criteria and alternatives grows. In large-scale portfolios, the cognitive burden of numerous pairwise comparisons can lead to inconsistencies, and the subjective nature of the inputs may introduce bias if not managed carefully [1].

Distance-based approaches, including the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) and the VIKOR method, are grounded in the geometric evaluation of alternatives relative to ideal and anti-ideal solutions [4]. They tend to perform well in environments where quantitative data are readily available and comparable across alternatives. For IT outsourcing companies that track detailed financial indicators, risk scores, and resource usage metrics, these

methods offer an efficient and reproducible way to produce rankings. Their outputs are straightforward to interpret: alternatives closer to the ideal point and further from the anti-ideal point are preferred. Nonetheless, these approaches require careful normalization of data, and their results can be sensitive to the weighting scheme and the presence of outliers [10].

Outranking methods such as ELECTRE and PROMETHEE adopt a different logic, focusing on the degree to which one alternative can be said to outrank another based on concordance and discordance measures [5, 7]. This family of methods is especially well suited to decision problems where both quantitative and qualitative criteria play decisive roles and where the objective is not merely to produce a strict ranking but to explore dominance relationships between alternatives. In IT outsourcing, such methods can be valuable when evaluating strategic partnerships or innovation-driven projects, where qualitative assessments – such as brand fit or potential for technology transfer – are as significant as measurable financial returns. However, their mathematical complexity and the need to calibrate thresholds or preference functions may pose challenges for organizations without specialized decision-analysis expertise.

Utility and efficiency-based methods, including Multi-Attribute Utility Theory (MAUT) and Data Envelopment Analysis (DEA), provide yet another perspective. MAUT aggregates multiple criteria into a single utility score, allowing explicit modelling of the decision-maker's risk preferences and trade-offs between attributes [8]. DEA, on the other hand, measures the relative efficiency of alternatives by comparing weighted outputs to weighted inputs [9]. These methods are particularly attractive to larger IT outsourcing firms with the capacity to collect and maintain high-quality datasets on project performance and resource utilization. In such cases, DEA can uncover underperforming projects or teams, while MAUT can integrate diverse performance dimensions into a unified decision framework. Their main limitation lies in the need for extensive, accurate, and comparable data, as well as in the technical expertise required to construct valid utility functions or linear programming models.

Hybrid methods, most notably AHP–TOPSIS, seek to combine the advantages of qualitative and quantitative approaches [3]. By using AHP to derive weights and TOPSIS to conduct the final ranking, these methods allow strategic priorities to be incorporated while ensuring that the evaluation reflects quantitative performance metrics. For IT outsourcing portfolios that include both strategic capability-building projects and high-revenue contracts, such hybrid models can offer a balanced view. They also facilitate structured participation from multiple

stakeholders, as the AHP stage can incorporate the perspectives of executives, project managers, and technical leads. The trade-off is that hybrid methods are more demanding in terms of modelling time and often require custom software or advanced spreadsheet models to implement effectively.

When comparing these approaches holistically, it becomes evident that no single method dominates across all decision contexts. Hierarchical methods excel in clarity and stakeholder engagement but may falter under the weight of large datasets. Distance-based methods are efficient for quantitative, well-structured data but require careful preprocessing and may oversimplify nuanced strategic considerations. Outranking methods handle complexity and mixed data types gracefully yet demand more advanced methodological skills. Utility and efficiency approaches provide a strong analytical foundation but depend heavily on data quality. Hybrid models offer the promise of integration but come with higher implementation costs.

Ultimately, the choice of method in IT outsourcing portfolio management must balance analytical rigour with practical constraints, aligning the method's strengths with the company's decision-making culture, data environment, and strategic objectives. The tables below synthesize these comparative insights, summarizing the advantages, limitations, and typical applications of each method within the IT outsourcing context

Table 2

Advantages, limitations, and IT outsourcing applications of MCDM methods

Method	Advantages	Limitations	IT outsourcing applications
AHP / ANP	Transparent; consistency check; intuitive for experts	Subjectivity; limited number of criteria	Determining weights for contract selection
TOPSIS / VIKOR	Quantitative accuracy; accounts for both benefits/costs	Sensitive to data normalization; full datasets needed	Ranking projects
ELECTRE / PROMETHEE	Works with incomplete/qualitative data	Complex interpretation; computationally demanding	Supplier or partner evaluation
MAUT / DEA	Integrated efficiency assessment	Requires large, accurate datasets	Evaluating resource efficiency
AHP–TOPSIS	Combines expert and quantitative evaluation	Requires specialized software	Comprehensive contract selection

Table 3

Computational characteristics of MCDM methods

Method	Complexity growth with alternatives	Sensitivity to weight changes	Need for normalization	Software availability
AHP / ANP	Quadratic (pairwise comparisons)	High	No	Expert Choice, Super Decisions
TOPSIS / VIKOR	Linear	Moderate	Yes	Excel, MATLAB, Python libs
ELECTRE / PROMETHEE	Quadratic	Low–Moderate	No	Decision Lab, Visual PROMETHEE
MAUT / DEA	Varies (linear; LP complexity)	Low	Sometimes	DEA Solver, R, Python
AHP–TOPSIS	Combination of above	Moderate–High	Yes	Custom integrations

3. Methodological Aspects of Selecting an MCDM Method

Selecting an appropriate multi-criteria decision-making method for IT outsourcing project portfolio management is not a purely technical exercise; it is a process shaped by the interplay of data availability, decision complexity, organizational priorities, and the cultural environment in which decisions are made. The methodological choice is ultimately a question of fit – between the capabilities of a given approach and the operational realities of the company.

One of the most decisive factors in this process is the nature and quality of the available data. When the evaluation must be based on qualitative judgments, incomplete information, or linguistic descriptors, hierarchical approaches such as the Analytic Hierarchy Process (AHP) or the Analytic Network Process (ANP) provide a structured way to translate expert opinions into numerical priorities [6]. By contrast, when companies can rely on complete and consistent quantitative indicators – such as net present value, risk exposure metrics, or resource allocation figures – distance-based techniques like TOPSIS and VIKOR can produce rankings that are both efficient to compute and easy to interpret [4]. In environments where data are mixed or inherently uncertain, fuzzy extensions of these methods offer an effective compromise, allowing qualitative assessments to be incorporated into quantitative models [3].

The size of the portfolio and the number of evaluation criteria also exert significant influence on method selection. Large portfolios containing dozens of alternatives can quickly render hierarchical methods impractical due to the exponential growth in pairwise comparisons. In such cases, approaches with linear computational complexity, including TOPSIS, VIKOR, or Data Envelopment Analysis (DEA), are generally more scalable [9]. Smaller portfolios with well-

defined strategic objectives, on the other hand, can benefit from the transparency and participatory nature of AHP, especially when stakeholder engagement is a priority.

Decision environments in IT outsourcing are often shaped by varying degrees of risk and uncertainty, particularly when projects involve new technologies, emerging markets, or untested business models. Under such conditions, outranking approaches like ELECTRE and PROMETHEE [5, 7] provide valuable flexibility by accommodating partial dominance and by allowing the decision-maker to set veto thresholds for certain criteria. Hybrid models that incorporate fuzzy logic can further enhance resilience to uncertainty by tolerating imprecision in both criteria weights and performance scores [3].

Practical considerations such as time constraints and resource availability frequently determine whether a method is viable in a given situation. When deadlines are tight and analytical resources are limited, simpler models such as TOPSIS or basic versions of the Multi-Attribute Utility Theory (MAUT) [8] can be implemented using standard spreadsheet tools, providing timely results without sacrificing basic methodological soundness. By contrast, strategic planning exercises that unfold over longer time horizons can justify the greater investment required to implement hybrid frameworks like AHP–TOPSIS, which combine the depth of qualitative weighting with the precision of quantitative ranking [3].

Finally, the growing availability of collaborative decision-support systems is changing the way MCDM methods are selected and applied in IT outsourcing contexts. Platforms such as *decideXpert* [3] integrate multiple decision-making algorithms into a single environment, enabling geographically distributed teams to participate in evaluations and receive real-time feedback on consistency and sensitivity. The integration of such tools not only streamlines the computational process but also improves transparency and stakeholder buy-in, reducing the likelihood of disputes over final rankings.

In practice, method selection is rarely about finding the perfect algorithm; rather, it is about aligning analytical capabilities with the strategic imperatives of the firm. Smaller organizations with limited data may value transparency and ease of use over methodological sophistication, whereas larger enterprises with mature analytics infrastructures can afford to deploy more complex, data-intensive models. The most effective choices are those that balance methodological rigour with usability, ensuring that portfolio decisions are not only technically defensible but also strategically coherent.

The key considerations outlined above can be synthesised into a condensed framework, presented in the table below, which maps different decision contexts to the most suitable MCDM methods for IT outsourcing portfolio management.

Conclusions

This extended review confirms that no single MCDM method universally fits all IT outsourcing portfolio management contexts. Hierarchical methods provide transparency and stakeholder engagement, distance-based methods offer computational efficiency, outranking methods handle uncertainty well, utility and efficiency methods excel with reliable data, and hybrid approaches balance strengths at the cost of complexity.

For practitioners, adopting a method selection framework such as Table 4 can align decision contexts with the most appropriate tools. For researchers, future work should explore integrating MCDM with AI-driven predictive analytics and digital twin simulations, enabling dynamic re-ranking as conditions evolve.

Table 4

Method Selection Map for IT Outsourcing Portfolios

Decision context	Data type	Portfolio size	Decision urgency	Recommended methods
Strategic, qualitative-heavy evaluation	Qualitative/mixed	Small	Low	AHP, ANP, Fuzzy AHP
Large portfolio, full metrics available	Quantitative	Large	Medium	TOPSIS, VIKOR, DEA
High uncertainty/risk	Mixed	Medium/Large	Medium	ELECTRE, PROMETHEE, Fuzzy MCDM
Rapid decision needed	Quantitative	Any	High	TOPSIS, MAUT
Balanced qualitative & quantitative	Mixed	Small/Medium	Low/Medium	Hybrid AHP–TOPSIS

References

1. M. C. Lacity, S. A. Khan, and L. P. Willcocks, «A review of the IT outsourcing literature: Insights for practice», *Journal of Strategic Information Systems*, vol. 18, no. 3, pp. 130–146, 2009. DOI: 10.1016/j.jsis.2009.06.002
2. Z. H. Wang, M. O. Esangbedo, and S. Bai, «Project portfolio selection based on multi-project synergy», *Journal of Industrial and Management Optimization*, vol. 19, no. 1, pp. 87–104, 2023. DOI: 10.3934/jimo.2021177
3. A. Saoud et al., «decideXpert: Collaborative system using AHP–TOPSIS and fuzzy techniques for multicriteria group decision-making», *SoftwareX*, vol. 25, p. 102026, 2025. DOI: 10.1016/j.softx.2024.102026
4. M. Behzadian, S. K. Otaghsara, M. Yazdani, and J. Ignatius, «A state-of-the-art survey of TOPSIS applications», *Expert Systems with Applications*, vol. 39, no. 17, pp. 13051–13069, 2012. DOI: 10.1016/j.eswa.2012.05.056

5. B. Roy, «The outranking approach and the foundations of ELECTRE methods», *Theory and Decision*, vol. 31, no. 1, pp. 49–73, 1991. DOI: 10.1007/BF00134132
6. T. L. Saaty, «Decision making with the analytic hierarchy process», *International Journal of Services Sciences*, vol. 1, no. 1, pp. 83–98, 2008. DOI: 10.1504/IJSSCI.2008.017590
7. J. P. Brans, P. Vincke, and B. Mareschal, «How to select and how to rank projects: The PROMETHEE method», *European Journal of Operational Research*, vol. 24, no. 2, pp. 228–238, 1986. DOI: 10.1016/0377-2217(86)90044-5
8. R. L. Keeney and H. Raiffa, *Decisions with Multiple Objectives: Preferences and Value Tradeoffs*. Cambridge, UK: Cambridge University Press, 1993. DOI: 10.1017/CBO9781139174084
9. W. D. Cook and L. M. Seiford, «Data envelopment analysis (DEA) – Thirty years on», *European Journal of Operational Research*, vol. 192, no. 1, pp. 1–17, 2009. DOI: 10.1016/j.ejor.2008.01.032
10. S. Opricovic and G. H. Tzeng, «Compromise solution by MCDM methods: A comparative analysis of VIKOR and TOPSIS», *European Journal of Operational Research*, vol. 156, no. 2, pp. 445–455, 2004. DOI: 10.1016/S0377-2217(03)00020-1
11. E. K. Zavadskas, Z. Turskis, and S. Kildienė, «State of art surveys of overviews on MCDM/MADM methods», *Technological and Economic Development of Economy*, vol. 20, no. 1, pp. 165–179, 2014. DOI: 10.3846/20294913.2014.892037
12. C. L. Hwang and K. Yoon, *Multiple Attribute Decision Making: Methods and Applications*. Berlin, Germany: Springer, 1981. DOI: 10.1007/978-3-642-48318-9

MODELING THE FORMATION OF DRONE SWARMS USING COMPONENT AND AGENT-BASED APPROACHES

Fedorovich O., Prohorov O., Leshenko Yu., Malieiev L. Kosenko V.

The task of studying the multi-component composition of military equipment (using UAVs as an example) during the delivery and formation of the necessary stocks of weapons and their components in a military conflict zone to ensure military superiority, which contributes to the successful achievement of military operation objectives, has been formulated and is being addressed. A multi-variant, multi-criteria problem is also being solved, which is related to the modeling of the logistics process of ensuring military superiority in a military conflict zone. The aim of the study is to model actions to provide combat units with various types of UAVs when forming a drone swarm to ensure military superiority in the air in conditions of threats and vulnerabilities associated with the logistics of supplying UAVs to the military conflict zone. A component approach is used to represent weapons and their components in the process of military equipment supply. A multi-layer component structure of UAVs is proposed for the systematic representation of weapons and their components in the process of supply and stockpiling in a military conflict zone. The factors of threats associated with the supply of the required types of UAVs in the quantities necessary to form a drone swarm are investigated using qualitative assessments by military experts and lexicographic ordering of possible UAV suppliers. The selection of UAV suppliers has been optimized using integer programming. A model has been created to optimize the logistics costs of supplying weapons and their components to the combat zone. Optimization tasks are solved under conditions of conflicting criteria of time, costs, and risks associated with the logistics of forming military superiority in the combat zone. An agent-based simulation model for researching the logistics of military cargo delivery in a conflict zone has been created. The results of the research can be used to model plans and logistics for the delivery of weapons and their components to combat zones in order to create military superiority.

Introduction

In modern conditions, the success of a military operation depends on many factors, including military superiority over a potential enemy [1, 2]. The use of military equipment in a state of war requires an analysis of the necessary quantity and range of weapons, taking into account the composition of their components, in order to create military superiority in the zone of military conflict [3, 4]. Creating superiority in weapons and military equipment is one of the important tasks in preparing for a military operation [5, 6].

Building military superiority is a complex logistical task that requires the involvement of developers, suppliers of weapons and military equipment, and participants in transport logistics in a heterogeneous transport environment

with possible transshipments from one transport route to another and temporary storage of military cargo [7].

There's some tricky logistics involved in delivering small batches of weapons and their components from manufacturers and suppliers who are far away from the combat zone [8–10]. When forming stocks of military equipment, it is necessary to take into account the composition of components that ensure the effective use of weapons (ammunition, spare parts, repair kits, auxiliary equipment) [11, 12]. The logistics of training military personnel to use modern weapons requires considerable time. All of the above, in general, has an impact on the formation of the combat capability of military units and the establishment of superiority over the enemy in the zone of military conflict [13].

Today, within the Unmanned Systems Forces (USF), strike companies of UAVs are being created, as well as other units specializing in various types of unmanned aerial vehicles. Providing the USF with the necessary unmanned aerial vehicles for forming drone swarms, equipment for their control, and ammunition is an important task.

Therefore, the topic of this publication is relevant, as it addresses and solves the problem of logistics for ensuring superiority over the enemy in a military conflict zone for the successful conduct of military operations under conditions of threats and vulnerabilities associated with the logistics of supplying weapons to UAV units [14, 15], as well as the use of the component method and a precedent-based approach to analyze the logistics of supplying and forming stocks of weapons and their components in the combat zone [16, 17].

1. Analysis of current publications and setting the task

An analysis of publications on this topic showed that existing works are mainly related to individual aspects of the task at hand, such as research on requirements for establishing military parity for the successful conduct of military operations [15], logistics of military cargo transportation [18], and uncertainties related to military transportation logistics [19]. There are no publications offering a comprehensive solution to the proposed task, analyzing threats related to the violation of military parity in a zone of military conflict, or logistical threats related to the transportation of military cargo to a combat zone. In addition, the impact of possible vulnerabilities related to the logistics of military cargo delivery to the combat zone, leading to a violation of military parity and, as a result, to the death of armed forces personnel, the destruction of defense infrastructure, the transition from offensive to defensive actions, etc., has not been studied.

The existing operational and tactical logistical problems related to establishing military superiority in a war zone have been analyzed, among which the following are noteworthy:

1. The problem of the large number of weapon components and small delivery batches to the war zone. In the existing literature, the task of forming stocks of weapons and their components is considered mainly in peacetime. Army stocks of weapons and manufacturers' stocks are formed, and long-term plans are implemented to form stocks of weapons for a sufficiently long period of time [19].

2. The logistical problem associated with the supply of weapons and their components through logistical channels in conditions of military threat. In peacetime, military equipment stocks are formed according to a pre-established schedule and logistical supply chains [18].

3. The problem of military threats, which leads to disruptions and stoppages in the supply of military equipment. In peacetime, threats are associated with possible transport breakdowns, accidents on transport routes, the influence of climatic factors, and the emergence of terrorist threats [20].

4. The problem of timely delivery of weapons and their components is related to military circumstances and the nature of combat operations in the zone of military conflict. In peacetime, when supplying military equipment, the requirements for the formation of stocks of weapons and their components are taken into account in accordance with pre-established plans for the formation of stocks in army warehouses [21].

The purpose of the study is to model actions to provide units with different types of UAVs when forming a drone swarm to ensure military superiority in the air in conditions of threats and vulnerabilities associated with the logistics of supplying UAVs to a military conflict zone using the component method to represent weapons and their components in the process of supplying military equipment.

The main criteria for evaluating the achievement of the research objective are threat levels, logistical risks, and the time and costs required to establish military superiority in a military conflict zone.

To achieve the research objective, the following tasks must be solved:

1. To form a component structure of military equipment for the systematic representation of weapons and their components (using UAVs as an example).
2. To develop a component method for creating innovative UAV models.
3. To analyze threat factors in the formation of military superiority in the air in a military conflict zone.

4. To conduct threat modeling in the logistics of production and supply of weapons and military equipment to ensure air superiority in a military conflict zone.

5. To optimize logistics costs for the formation of stocks of weapons and their components in a military conflict zone.

6. To develop a method for improving the efficiency of weapons use by creating asymmetry in military parity in a military conflict zone.

7. To create an agent-based simulation model for researching the logistics of production and supply of weapons and military equipment to a military conflict zone to establish military superiority.

2. Modeling the provision of UAV units in a military conflict zone using a component-based approach

2.1 Formation of the UAV component structure for systematic representation of weapons and their components

For the effective use of weapons in a military conflict zone, it is necessary to comprehensively supply, in the required quantities, all components of military equipment (ME) related to the effective functioning of weapons:

- basic structural components (BSC);
- ammunition (A);
- spare parts and repair kits (SP);
- auxiliary equipment (AE), which for UAVs includes control, communication, power supply, transportation, and training equipment;
- mobile workshops for repairing military equipment (MME).

The lack of the necessary quantity of weapon components in the war zone (in accordance with NATO requirements and standards) leads to a reduction in the combat effectiveness of military equipment, as well as the inefficient use of weapons. Therefore, it is important to systematically present the architecture of military equipment in the form of weapons and their supporting components.

Fig. 1 shows a hierarchical diagram of types of weapons, highlighting the component structure of a particular type of military equipment.

Here, on the top level, you can see the types of weapons used in war zones, with aviation equipment, including UAVs, highlighted on the second level. The third level shows the components.

For example, a typical copter-type UAV consists of the following components:

- body or fuselage (B);
- power unit (engine) (PU) ;
- autopilot (flight controller) with sensors (AP);

- control equipment (CE);
- air screw (propeller) (AS);
- information signal and telemetry transmitter and receiver (ITR);
- onboard electronics (OE);
- power source (battery, generator) (PS);
- landing gear (chassis) (LG).

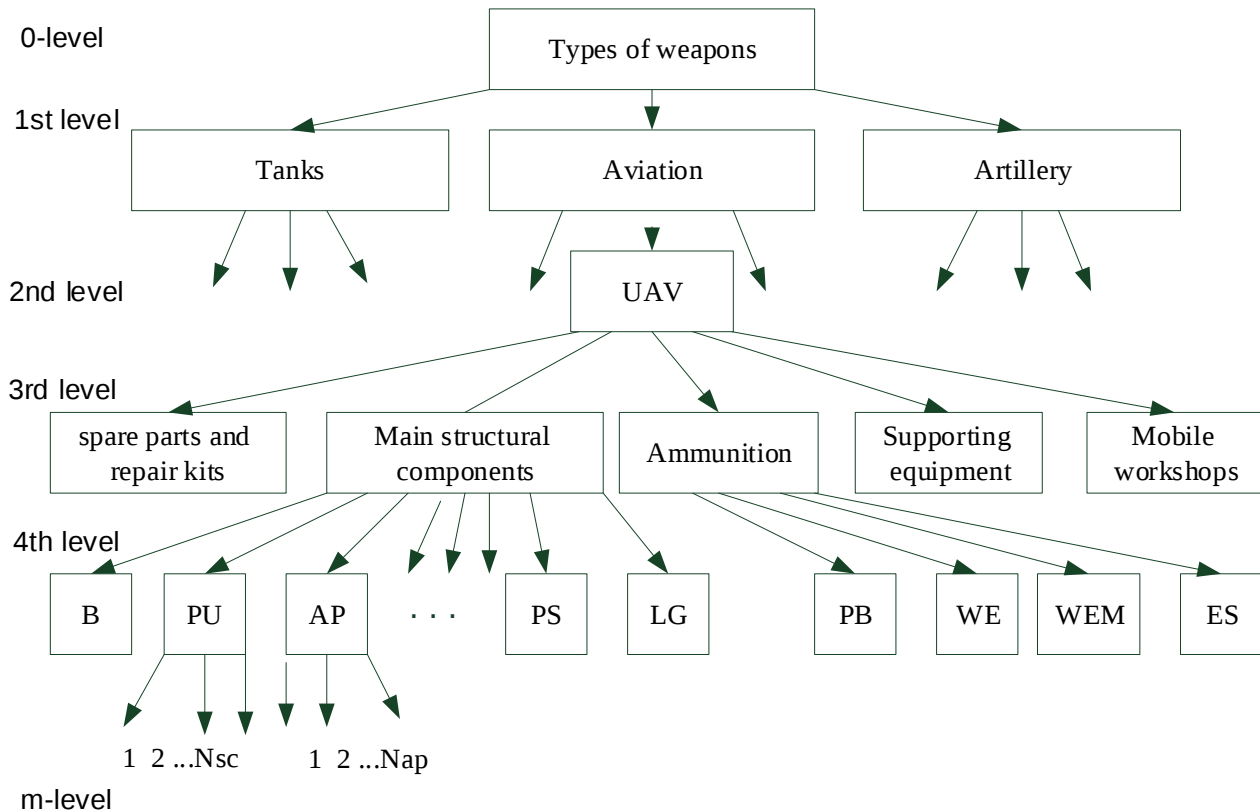


Fig. 1. Component structure of a type of military equipment (UAV)

At the next level are components that make up the supporting elements. For strike UAVs, ammunition is an important component. Combined-type UAVs are adapted to drop combat grenades (F-1, RGD-5, RGO, RGN, RK-3) and mines (VOG-17, 82 mm for the 2B14 mortar and 120 mm for the 2B11 and 2B12) on the enemy. UAVs are transformed into high-precision weapons that simultaneously perform aerial reconnaissance tasks. In general, ammunition contains the following components:

- projectile body (PB);
- warhead elements (cone or arrow) (WE),
- warhead element mountings (WEM),
- explosive substance (ES).

The details of military equipment can be further elaborated, taking into account not only the processes of supplying weapons and their components to the zone of military conflict, but also the processes of manufacturing military equipment.

The modern approach to the creation of military equipment, based on NATO standards and regulations, allows new weapons to be created by combining existing components (reusable components). Therefore, it is advisable to use the component method when modeling the creation of new types of weapons.

The following typical components can be identified for the creation of new military equipment:

- reusable components (RC), which can be used in various types and kinds of weapons;
- components that, through the adaptation of RC (AC), can be used in the creation of new types of weapons;
- new components (IC) of an innovative nature. These components must be used to create innovative models of military equipment.

The creation of innovative types of UAVs with the predominant use of RCC allows minimizing risks, development time, and costs. However, the tactical and technical characteristics of UAVs and their combat effectiveness will not differ significantly from existing ones. The predominant use of innovative components allows for improving the tactical and technical characteristics for combat use of UAVs (accuracy, range, kill zone, etc.), but at the same time increases the risks, time, and costs of an innovative project. A compromise lies in the combined use of RC, AC, and NC components in the military equipment being developed. This allows for the creation of innovative UAV models with improved characteristics while keeping project time and cost risks at a satisfactory level.

The use of the component method to represent the multilayered structure of military equipment allows for a systematic analysis of the logistics of supplying and forming comprehensive stocks of UAVs and their components in a military conflict zone, taking into account two main areas [20]:

1. Supply of UAVs from army warehouses and manufacturers' warehouses.
2. Production of UAVs for delivery and formation of necessary stocks in the zone of military conflict.

The logistics of forming stocks of military equipment in the zone of military conflict belongs to the «pull» type of logistics system. In this case, the requirements for UAV stockpiles are determined at the operational-tactical level, taking into account the situation and objectives of military operations in the war zone. These requirements are initially analyzed at the army depot level. If army depots

do not have sufficient stockpiles, the possibility of supplying them from manufacturers' warehouses is analyzed. If the requirements for the comprehensive formation of UAV stocks in the zone of military conflict cannot be met using army warehouses and manufacturers' warehouses, production plans are drawn up to eliminate the shortage of UAVs and their components.

The component method allows, taking into account the use of RC, AC, and NC components, to determine the logistics requirements for the comprehensive formation of military equipment stocks in a military conflict zone in the following sequence of actions:

1. First, at the operational and tactical level, using expert assessments of military specialists, requirements are formed for the required number W_j of the j -th type of UAVs for use in the zone of military conflict. At the first level, requirements for the supporting components of the weapon W_{j1} are formed.

2. The availability of the required number of certain types of UAVs W_{j1} in army warehouses is checked. If

$$W_{j1} \leq W'_{j1} \quad (1)$$

where W'_{j1} is the availability of UAVs of the j -th type in the army warehouses), then the required amount of W_{j1} is supplied to the zone of military conflict to form stocks of the j -th type of UAV.

3. If

$$W_{j1} > W'_{j1}, \quad (2)$$

then the availability of W'_{j1} in the manufacturers' warehouses is checked.

If

$$W_{j1} \leq W'_{j1} + W''_{j1}, \quad (3)$$

then W_{j1} is delivered to the military conflict zone.

4. If

$$W_{j1} > W'_{j1} + W''_{j1}, \quad (4)$$

then in this case production plans are formed for the output

$$\Delta W_{j1} \geq W_{j1} - (W'_{j1} + W''_{j1}). \quad (5)$$

5. To implement the production plans ΔW_{j1} , it is necessary to form auxiliary plans for the production of supporting components of the third level of representation of the component architecture of military equipment (see Fig. 1).

Let's introduce the following notation:

B – ammunition;
 Z – spare parts;
 V – auxiliary equipment;
 M – mobile workshops.

Given these notations, we write the following constraints:

$$\begin{aligned}
 \Delta W_{j2_B} &\geq W_{j2_B} - (W'_{j2_B} + W''_{j2_B}), \\
 \Delta W_{j2_Z} &\geq W_{j2_Z} - (W'_{j2_Z} + W''_{j2_Z}), \\
 \Delta W_{j2_V} &\geq W_{j2_V} - (W'_{j2_V} + W''_{j2_V}), \\
 \Delta W_{j2_M} &\geq W_{j2_M} - (W'_{j2_M} + W''_{j2_M}),
 \end{aligned} \tag{6}$$

where $\Delta W_{j2_B}, \Delta W_{j2_Z}, \Delta W_{j2_V}, \Delta W_{j2_M}$ is the required amount of ammunition, spare parts, auxiliary equipment, and mobile workshops for the comprehensive supply of UAVs to the military conflict zone.

6. Further, if there is a need to produce components for the i -th level of representation of the component architecture of military equipment, then in this case it is necessary:

$$\begin{aligned}
 \Delta W_{ji_B} &\geq W_{ji_B} - (W'_{ji_B} + W''_{ji_B}), \\
 \Delta W_{ji_Z} &\geq W_{ji_Z} - (W'_{ji_Z} + W''_{ji_Z}), \\
 \Delta W_{ji_V} &\geq W_{ji_V} - (W'_{ji_V} + W''_{ji_V}), \\
 \Delta W_{ji_M} &\geq W_{ji_M} - (W'_{ji_M} + W''_{ji_M}).
 \end{aligned} \tag{7}$$

Thus, the use of the component method in the integrated formation of UAV stocks of various types, using warehouses and production of components, allows to formulate requirements for the volume of stocks and production plans for different levels of component representation.

2.2 Component method for creating innovative UAV prototypes

When creating new models of military equipment, not only reusable components (RC) but also adaptive components (AC) and innovative (IC) components may be used in the product. At the initial stages of military equipment development, it is necessary to form the component architecture of the product, taking into account the tactical and technical characteristics (TTC) set out in the project specifications.

A method for synthesizing innovative UAV models has been proposed, taking into account multi-level component architecture, which is based on the active use of existing components (EC), allowing to minimize project risks and reduce the time required to create relevant weapons for use in a military conflict zone. The component method consists of the following stages:

1. At the beginning, the component composition of the ME product is formed at the subsystem level. TTCs are specified for each component. We search for existing components with the required TTCs in a database formed based on previous developments and existing weapons. The component database will be presented in the form of precedent bases (PB), which accumulate the experience of previous developments.

2. During the synthesis process, the following directions for creating the component architecture of innovative UAV models are possible:

A) the PB database contains the necessary components for forming the component architecture of the product at the first level (subsystem level). In this case, we proceed to the next stage of design;

B) the precedent database contains components whose TTC does not fully match those required in the technical specifications for the project. In this case, it is necessary to search for similar components by TTC value. That is, it is necessary to evaluate the similarity between the k -th precedent in the PB and the required p -th component of the project. In doing so, both quantitative and qualitative metrics can be used to evaluate similarity. In the simplest method of evaluating similarity (Minkowski metric), Euclidean distance is used:

$$\lambda_{kp} = \sqrt{\sum_{l=1}^L \rho_l (\lambda_{kl} - \lambda_{pl})^2}, \quad (8)$$

where $l = \overline{1, L}$ – 1st characteristic from TTC components in the project;

λ_{pl} – value of the 1st characteristic for the r -th UAV component in the project;

λ_{kl} – value of the 1st characteristic for the k -th precedent in PB;

ρ_l – significance (weight) of the 1st characteristic.

In the early stages of creating a new military product, there are characteristics whose proximity to the project can be assessed qualitatively. In this case, the proximity assessment can be carried out using a qualitative representation.

For example, in the form of letters of the Latin alphabet:

$$\partial_{kp_l} = \begin{cases} A - \text{maximum proximity of components} \\ \quad \text{according to the 1st characteristic;} \\ B - \text{satisfactory proximity of components} \\ \quad \text{according to the 1st characteristic;} \\ C - \text{allowable proximity of components} \\ \quad \text{according to the 1st characteristic;} \\ D - \text{remoteness of components} \\ \quad \text{according to the 1st characteristic.} \end{cases} \quad (9)$$

Next, depending on the importance of the TTC UAV, an ordered list of characteristic values (in the form of «words») is compiled. For example: B, B, A, C, B for five characteristics. Then, if there is not only one possible close component available, but several in the precedent database, a list of «words» is formed. For example:

B, B, C, A, B
B, B, D, C, B
C, B, D, C, B
B, B, A, C, B
A, C, D, A, C.

The list of «words» is lexicographically sorted and takes the following form:

A, C, D, A, C
B, B, A, C, B
B, B, C, A, B
B, B, D, C, B
C, B, D, C, B.

Next, the permissible proximity indication is set in the form of a «control word». For example:

$\boxed{B, B, B, B, B}$.

The «control word» is placed in an ordered list of «words»:

A, C, D, A, C
B, B, A, C, B
 $\boxed{B, B, B, B, B}$
B, B, C, A, B
B, B, D, C, B
C, B, D, C, B.

Next, select the closest component by the characteristic value in PB, which is adjacent to the «control word»: B, B, A, C, B. This component will be used in further design;

C) there is no component in the PB of existing components with the required proximity in terms of characteristics. In this case, we proceed to the decomposition of the i -th level project component into possible components of the lower $(i+1)$ level of the multi-component architecture of the product and search for the required components in the PB at the $(i+1)$ level of representation. Next, we repeat the actions in step 2;

D) a possible situation is related to the design of a new innovative component. Figure 2 shows a diagram of the component synthesis of military equipment.

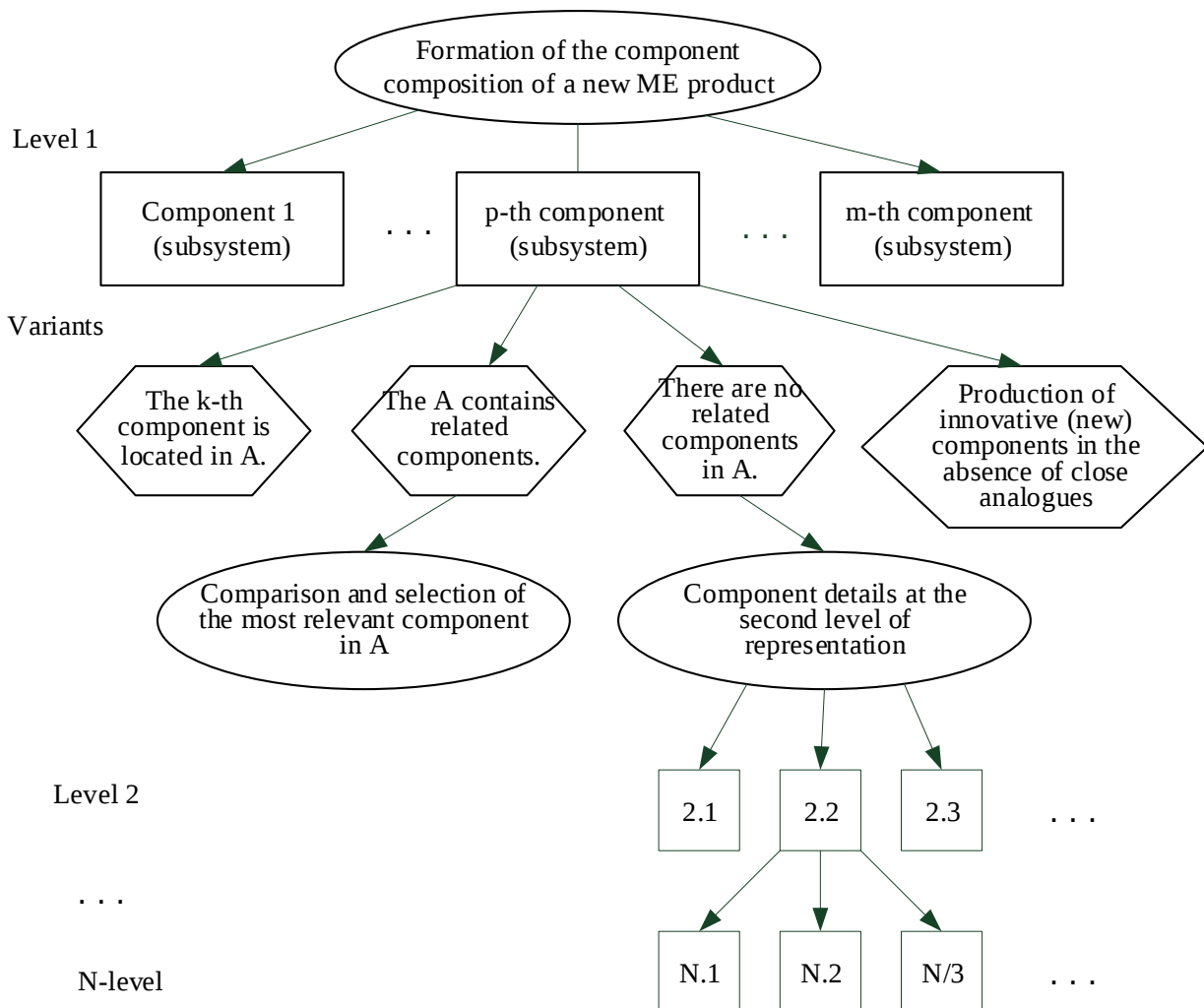


Fig. 2. Component synthesis diagram for a new military equipment product

2.3. Optimization of logistics costs for the supply of weapons and components in the combat zone

The complexity of the integrated supply of weapons to the combat zone is associated with the many supporting components of military equipment that are necessary for the effective use of weapons (ammunition, spare parts, repair kits, auxiliary equipment for maintenance, training of military personnel, etc.) Therefore, the task of optimizing logistics costs for the timely delivery of UAVs and their components to the combat zone to create a military advantage is urgent. The effectiveness of the use of a particular j -th type of UAV in the combat zone will depend on the comprehensive supply of the required number of supporting components (designated respectively as B, Z, V, M). To create a stockpile of UAVs of the j -th type in the amount of W_j , it is necessary that the number of UAVs is in the interval:

$$W_{j\min} \leq W_j \leq W_{j\max}, \quad (10)$$

here $W_{j\min}$ corresponds to the insurance stocks W_j in the combat zone, which ensure, taking into account military risks, the establishment of parity of forces for the j -th type of A (without violating the nature of hostilities).

The value of $W_{j\max}$ creates an asymmetry in the military parity of forces, which can ensure a change in the nature of hostilities (transition from defensive to offensive actions). Therefore, planning of weapons stocks within the limits of (10) is the basis for the successful fulfillment of the objectives of a military operation. Also, it is necessary to fulfill similar requirements for the formation of stocks of supporting components of the j -th type of UAV (B, Z, V, M):

$$\begin{aligned} W_{jB\min} &\leq W_{jB} \leq W_{jB\max}, \\ W_{jZ\min} &\leq W_{jZ} \leq W_{jZ\max}, \\ W_{jV\min} &\leq W_{jV} \leq W_{jV\max}, \\ W_{jM\min} &\leq W_{jM} \leq W_{jM\max}. \end{aligned} \quad (11)$$

The complex supply of UAV types and its supporting components can be carried out through a variety of logistics channels (LC). Therefore, it is relevant to solve the problem of optimizing the choice of LCs, taking into account logistics costs and requirements for the timely formation of military equipment stocks in the CZ, under conditions of possible military risks R_j . The optimization problem belongs to the class of multivariate problems using a set of indicators and constraints. To solve this optimization problem, we will use the method

of integer (Boolean) programming. Let's introduce a Boolean variable x_{jik} , which has the following values:

$$x_{jik} = \begin{cases} 1, & \text{if } k\text{-th logistics channel is selected} \\ & \text{for the supply of } i\text{-th components of } j\text{-th type of weapons;} \\ 0, & \text{other wise.} \end{cases} \quad (12)$$

Let us present the logistics indicators of military equipment supply in the combat zone as follows:

1. The volume of UAV stocks that are formed in the area of military conflict:

$$W = \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} w_{jik} x_{jik}, \quad (13)$$

where w_{jik} is the number of i -th components of the j -th type of UAV, which are supplied according to the k -th possible LC;

n_{ij} is the number of possible LCs from which the i -th components of the j -th UAV type are selected for supply;

n_j is the set of different components of the j -th UAV type;

Q is the number of UAV types entering the military conflict zone;

$i = 1$ corresponds to the weapon, $i = 2, \dots$, refers to the supporting components of the weapon.

2. Time of delivery of UAVs to the zone of military conflict:

$$T = \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} t_{jik} x_{jik}, \quad (14)$$

where $T \leq T'$, T' is the permissible time that ensures timely delivery of weapons to the military conflict zone;

t_{jik} is the time of delivery of batches of the i -th component for the j -th type of weapon according to the selected k -th LC.

3. Risks of supplying UAVs to a military conflict zone under martial law:

$$R = \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} r_{jik} x_{jik}, \quad (15)$$

where $R \leq R'$, R' is the permissible risk of supplying weapons to the CZ;

r_{jik} is the risk of supplying batches of i -th components for the j -th type of weapon according to the selected k -th LC.

4. Logistical costs associated with the supply of UAVs to the zone of military conflict:

$$L = \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} l_{jik} x_{jik}, \quad (16)$$

where $L \leq L'$, L' is the allowable logistics costs of supply,

l_{jik} is the logistics costs associated with the supply of the i -th component for the j -th type of UAV according to the selected k -th LC.

To form the necessary UAV stocks in a military conflict zone under conditions of military risks, in order to create a military advantage, it is necessary to:

$$\max W, W = \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} w_{jik} x_{jik}, \quad (17)$$

taking into account

$$W_{jik_{\min}} \leq W_{jik} \leq W_{jik_{\max}} \quad \text{for all } j, i, k. \quad (18)$$

The following conditions must be met:

$$\begin{aligned} T \leq T', T &= \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} t_{jik} x_{jik}, \\ R \leq R', R &= \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} r_{jik} x_{jik}, \\ L \leq L', L &= \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} l_{jik} x_{jik}. \end{aligned} \quad (19)$$

For the timely supply and formation of the necessary stocks of UAVs in the zone of military conflict, under martial law, it is necessary:

$$\min T, T = \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} t_{jik} x_{jik}, \quad (20)$$

Subject to the following restrictions:

$$\begin{aligned} W \geq W', W &= \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} w_{jik} x_{jik}, \\ R \leq R', R &= \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} r_{jik} x_{jik}, \\ L \leq L', L &= \sum_{j=1}^Q \sum_{i=1}^{n_j} \sum_{k=1}^{n_{ij}} l_{jik} x_{jik}. \end{aligned} \quad (21)$$

The solution of the set optimization problems with Boolean variables can be carried out by one of the possible methods of integer optimization (for example, the modified branch-and-bound method).

3. Modeling the logistics of establishing military superiority in a combat zone

3.1 Analysis of Threat Factors in Establishing Military Superiority in the Combat Zone

In order to establish military superiority, it is necessary to analyze the main components of weapons and military equipment in a particular area of military conflict. For example, such components in the form of threat factors are the lack of the following types of weapons in the required quantity:

- aviation (including unmanned aerial vehicles)
- heavy weapons,
- ammunition.

Suppose that, with the help of military experts, the required types and number of UAVs for the area of military conflict under consideration are determined in the form of a set – M , $m_i \in M$, $i = \overline{1, N}$. Where m_i is the number of UAVs of the i -th type required to establish military superiority, N is the number of required types of UAVs to establish military superiority in the combat zone (CZ).

It is known, based on the available weapons, how many UAVs of the i -th type are available in the CZ – $m_{0i} \in M_0$, where m_{0i} is the availability of weapons of the i -th type at the present time.

Then

$$\Delta m_i = m_i - m_{0i} \quad (22)$$

is the number of weapons of the i -th type required to ensure military parity of forces in the CZ. The greater is the value of Δm_i , the greater is the threat associated with the lack of the required number of weapons of the i -th type. The emergence of threats associated with the required number of certain types of UAVs can lead to losses in the CZ. Therefore, it is necessary to assess the level of threats, taking into account possible factors related to the insufficient number of weapons.

For the convenience and simplicity of analyzing the level of threats, we use the estimates of military experts for each i -th type of UAV.

To do this, we will use qualitative assessments of the level of threats in the form of linguistic variable values [21]:

$$x_i = \begin{cases} A - \text{low threat level} \\ \quad \text{to the } i\text{-th type of weapon;} \\ B - \text{acceptable level of threat;} \\ C - \text{high level of threat.} \end{cases} \quad (23)$$

For each i -th type of UAV, there is a set of suppliers P_i that can supply m_{ij} weapons $m_{ij} \leq \Delta m_i$. The value of m_{ij} is a factor that affects the threat level x_i and is determined with the help of military experts for each i -th type of weapon. In addition, it is necessary to take into account the time of production and delivery t_{ij} of the i -th type of UAV by the j -th supplier. Let w_{ij} denote the costs of production and delivery of the i -th type of UAV by the j -th supplier, and r_{ij} denote the risks of production and delivery of the i -th type of UAV by the j -th supplier.

Let's define these values in a qualitative scale as multiples of three estimates:

$$\begin{aligned} t_{ij} &= \begin{cases} A - \text{minimum delivery term;} \\ B - \text{acceptable delivery term;} \\ C - \text{maximum delivery term;} \end{cases} \\ w_{ij} &= \begin{cases} A - \text{minimum costs;} \\ B - \text{acceptable costs;} \\ C - \text{maximum costs;} \end{cases} \\ r_{ij} &= \begin{cases} A - \text{minimum risk;} \\ B - \text{acceptable risk;} \\ C - \text{maximum risk.} \end{cases} \end{aligned} \quad (24)$$

To determine the supplier in the set P_i , we represent each possible j -th supplier for the i -th type of weapon as a sequence of qualitative values x_{ij} , t_{ij} , w_{ij} , r_{ij} , where x_{ij} is the level of threat established by military experts depending on the value m_{ij} , which is associated with the number of the i -th type of UAV and can be supplied by the j -th supplier.

The list of qualitative values x_{ij} , t_{ij} , w_{ij} , r_{ij} will be represented as a «word». Then, for the set P_i of possible suppliers of the i -th type of UAV, we will get a list of «words» that will represent the options chosen by the suppliers.

For example, we have 10 options:

1. A, C, B, B
2. B, A, B, B
3. B, B, A, B
4. C, A, A, B
5. B, B, C, B
6. B, A, C, C
7. C, A, A, A
8. B, B, A, C
9. C, A, A, C
10. A, B, B, C.

By ordering the options in the lexical and graphical sense, we get:

10. A, B, B, C
1. A, C, B, B
2. B, A, B, B
6. B, A, C, C
3. B, B, A, B
8. B, B, A, C
5. B, B, C, B
7. C, A, A, A
4. C, A, A, B
9. C, A, A, C.

It is noticeable that the options with the lowest level of threat will be in the top half of the ordered list of possible suppliers of UAVs of the i -th type. In this example, it is advisable to take the second option of the supplier with an acceptable level of threats, the minimum period of production and delivery of the i -th type of UAV, the acceptable value of production and delivery costs, and the acceptable value of risk.

Thus, for each i -th type of UAV, suppliers will be selected that ensure the minimum level of threat for all weapons in the formation of military parity of forces in the area of military conflict.

Given the large number of possible UAV suppliers, we will use quantitative estimates to formulate an optimization problem of selecting a supplier from a set of suppliers that minimizes the level of threats in the formation of air superiority in the combat zone.

Let z_{ij} be a Boolean variable that takes the values:

$$z_{ij} = \begin{cases} 1 - \text{if the } i\text{-th supplier is selected} \\ \text{for the } j\text{-th type of weapons;} \\ 0 - \text{other wise;} \end{cases} \quad (25)$$

wherein

$$\sum_{j=1}^{n_i} z_{ij} = 1. \quad (26)$$

Then, the threat criterion can be given in the following form:

$$V = \sum_{i=1}^N \sum_{j=1}^{n_i} v_{ij} z_{ij}, \quad (27)$$

where v_{ij} is a quantitative assessment of the level of threat proposed by military experts (for example, on a 10-point scale) when choosing the j -th supplier of the i -th type of UAV,

n_i is the number of possible suppliers of the i -th type of UAV.

It should be noted that v_{ij} depends on the number of $m'_{ij} \leq \Delta m_{ij}$, where m'_{ij} is the actual number of UAVs of the i -th type that can be supplied by the j -th supplier.

The total time for the production and delivery of weapons (a pessimistic estimate of the time associated with the sequential nature of the supply) is determined by the formula:

$$T = \sum_{i=1}^N \sum_{j=1}^{n_i} t_{ij} z_{ij}. \quad (28)$$

The costs associated with the production and supply of the necessary weapons and military equipment to establish military superiority in the CZ are:

$$W = \sum_{i=1}^N \sum_{j=1}^{n_i} w_{ij} z_{ij}. \quad (29)$$

Risks of production and supply of weapons to establish military superiority in the CZ are determined by the formula:

$$R = \sum_{i=1}^N \sum_{j=1}^{n_i} r_{ij} z_{ij}. \quad (30)$$

It is necessary to minimize the threat associated with the lack of the required number of UAVs of various types in the CZ:

$$\min V, V = \sum_{i=1}^N \sum_{j=1}^{n_i} v_{ij} z_{ij}, \quad (31)$$

if the following restrictions are met:

$$\begin{aligned} T &\leq T', T = \sum_{i=1}^N \sum_{j=1}^{n_i} t_{ij} z_{ij}, \\ W &\leq W', W = \sum_{i=1}^N \sum_{j=1}^{n_i} w_{ij} z_{ij}, \\ R &\leq R', R = \sum_{i=1}^N \sum_{j=1}^{n_i} r_{ij} z_{ij}, \end{aligned} \quad (32)$$

where T', W', R' – respectively, the permissible values of the time of production and delivery of UAVs, permissible costs and risks of establishing a military advantage in the CZ.

In the case of possible parallel deliveries of UAVs to the CZ (optimistic estimate of the delivery time), the time is defined as:

$$T = \max_{i=1}^N \left(\sum_{j=1}^{n_i} t_{ij} z_{ij} \right). \quad (33)$$

If military experts have difficulty determining threat assessments, then in this case we will use the values

$$\Delta m'_{ij} = \Delta m_{ij} - m'_{ij}, \quad (34)$$

which must be minimized during the production and delivery of UAVs to the CZ. In this case, it is necessary:

$$\min \Delta m', \Delta m' = \sum_{i=1}^N \sum_{j=1}^{n_i} \Delta m'_{ij} z_{ij}, \quad (35)$$

taking into account the permissible values;

$$\begin{aligned} T &\leq T', T = \sum_{i=1}^N \sum_{j=1}^{n_i} t_{ij} z_{ij}, \\ W &\leq W', W = \sum_{i=1}^N \sum_{j=1}^{n_i} w_{ij} z_{ij}, \\ R &\leq R', R = \sum_{i=1}^N \sum_{j=1}^{n_i} r_{ij} z_{ij}. \end{aligned} \quad (36)$$

If, for the i -th type of UAV, it is possible to use not one supplier, but a group of suppliers, then in this case, it is necessary to evaluate each k -th group of possible suppliers for the j -th option of selecting suppliers of the i -th type of weapon:

$$\begin{aligned} v_{ij} &= \sum_{k=1}^{l_j} v_{ijk}, \\ t_{ij} &= \sum_{k=1}^{l_j} t_{ijk}, \\ w_{ij} &= \sum_{k=1}^{l_j} w_{ijk}, \\ r_{ij} &= \sum_{k=1}^{l_j} r_{ijk}, \end{aligned} \tag{37}$$

where l_j is the number of possible suppliers in the group for the j -th option of selecting possible suppliers.

Then, to assess the level of threats, it is necessary to minimize the level of threat:

$$\min V, \quad V = \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} v_{ijk} \right) z_{ij}, \tag{38}$$

while fulfilling the constraints:

$$\begin{aligned} T &\leq T', \quad T = \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} t_{ijk} \right) z_{ij}, \\ W &\leq W', \quad W = \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} w_{ijk} \right) z_{ij}, \\ R &\leq R', \quad R = \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} r_{ijk} \right) z_{ij}. \end{aligned} \tag{39}$$

If there are difficulties associated with quantifying the military threat level v_{ijk} , then it is necessary to:

$$\min \Delta m', \quad \Delta m' = \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} \Delta m'_{ijk} \right) z_{ij}, \tag{40}$$

taking into account the fulfillment of the constraints

$$\begin{aligned}
T \leq T', \quad T &= \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} t_{ijk} \right) z_{ij}, \\
W \leq W', \quad W &= \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} w_{ijk} \right) z_{ij}, \\
R \leq R', \quad R &= \sum_{i=1}^N \sum_{j=1}^{n_i} \left(\sum_{k=1}^{l_j} r_{ijk} \right) z_{ij}.
\end{aligned} \tag{41}$$

3.2 Modeling threats in the logistics of UAV production and supply

In wartime, the logistics chains used for the production and supply of weapons and military equipment to the conflict zone may be disrupted by threats related to the aggressor's actions. At the same time, the realization of threats leads to the excitation of various vulnerabilities in the logistics of production and supply, which can ultimately lead to losses in the area of military conflict. The logical chain that needs to be investigated is: threat – vulnerabilities - losses.

Vulnerabilities may be related to the condition and bottlenecks in the heterogeneous transportation network used to transport weapons and military equipment. Threats to vulnerabilities may include:

- moral and physical obsolescence of the transportation system,
- bottlenecks in the form of transshipment of military cargo from one transportation highway to another,
- necessary warehousing and temporary storage of military cargo in certain locations of the transportation network,
- transportation of ammunition that may explode due to the actions of the aggressor,
- violation of cargo size and weight requirements,
- violation of special transportation regime requirements, etc.

In the event of a threat from the aggressor (air raids, missile attacks, long-range artillery, etc.), vulnerabilities are triggered, which lead to damage in the area of military conflict. Therefore, there is an urgent task to determine the impact of threats on the exploitation of vulnerabilities that lead to losses. This affects the nature of hostilities in the zone of military conflict.

To analyze the logical sequence: threat – vulnerabilities – losses, we will use the estimates of military experts, which are formed using a full factorial experiment (FFE). As an example of such an analysis, let's consider the impact

of a threat in the form of an air raid on a section of a railroad on which a batch of UAVs is being transported. Let's assume that military experts who know the specific route of the military cargo have identified three possible transportation vulnerabilities in the form of factors x_j of the FFE:

- cargo transshipment (x_1),
- temporary storage (x_2),
- violation of the special transportation regime (x_3).

Table 1 shows the FFE plan, where the plan lines correspond to possible vulnerabilities (no vulnerabilities – the first line $(-1, -1, -1)$, the presence of all vulnerabilities – the last line of the plan $(+1, +1, +1)$).

Table 1

FFE plan for loss assessment

	x_1	x_2	x_3	y
1	-1	-1	-1	0
2	-1	-1	+1	5
3	-1	+1	-1	3
4	-1	+1	+1	8
5	+1	-1	-1	2
6	+1	-1	+1	7
7	+1	+1	-1	5
8	+1	+1	+1	10

The right column of the FFE plan contains military experts' estimates of the levels of damage that would result from the vulnerabilities in the event of an air attack, for example, on a ten-point scale. The FFE plan can be used to build a regression relationship that allows you to estimate the impact of individual factors on damage, provided that certain vulnerabilities are triggered in the event of a threat (air raid):

$$y = b_0 + b_1x_1 + b_2x_2 + b_3x_3 + b_{12}x_1x_2 + b_{13}x_1x_3 + b_{23}x_2x_3 + b_{123}x_1x_2x_3, \quad (42)$$

where $b_0, b_1, b_2, b_3, b_{12}, b_{13}, b_{23}, b_{123}$ are the coefficients associated with the influence of factors x_1, x_2, x_3 on the value of the amount of damage (y).

Let's build the linear part of the regression dependence associated with the impact of individual factors (vulnerabilities) on the amount of damage.

After simple calculations, we get:

$$y = b_0 + b_1x_1 + b_2x_2 + b_3x_3 = 5 + x_1 + 1,5x_2 + 2,5x_3. \quad (43)$$

The dependence obtained as a result of the virtual experiment indicates the influence of individual factors:

1. According to experts, the most important vulnerability that is triggered by an air raid is the violation of the special regime for the transportation of military cargo ($b_3 = 2,5$).

2. Less important is the vulnerability associated with the temporary storage of military cargo ($b_2 = 1,5$).

3. Vulnerability associated with the transshipment of military cargo is the least important in terms of the impact on damage in the event of an air raid ($b_1 = 1$).

3.3 A method of increasing the effectiveness of the use of weapons by creating an asymmetry in military parity of forces in a military conflict zone

Military parity of forces is associated with the use of almost identical types of enemy weapons in a zone of military conflict. The emergence of new types of weapons and military equipment makes it possible to use more effective weapons in the conflict zone that far exceed the combat characteristics of their analogues (for example, the use of a swarm of drones with increased range and accuracy of impact). Therefore, by creating an asymmetry in the number of weapons and the use of different types of UAVs in a swarm, it is possible to increase the overall effectiveness of UAVs in a military conflict zone. Therefore, it is important to develop a method to increase the effectiveness of the use of weapons in a zone of military conflict by creating an asymmetry in the types of weapons in the military parity of forces. Due to the multivariate nature of the asymmetry in the military parity of forces, we will use the method of integer (Boolean) programming as a mathematical tool for solving the problem.

Let x_{ijk} be a Boolean variable with the following values:

$$x_{ijk} = \begin{cases} x_{ijk} = 1, & \text{if the } j\text{-th type of weapon} \\ & \text{and its } k\text{-th supplier are selected} \\ & \text{for the } i\text{-th type of weapons;} \\ x_{ijk} = 0, & \text{other wise.} \end{cases} \quad (44)$$

Further, with the help of military specialists (experts), for each type of UAV, the effectiveness q_{ij} of using UAVs in a military conflict zone is determined.

Then, for each type of UAV, the efficiency of use is as follows:

$$Q_i = \sum_{j=1}^{n_i} \sum_{k=1}^{p_j} q_{ij} m_{ijk} x_{ijk}, \quad (45)$$

where m_{ijk} is the number of units of weapon components (ammunition, components, supporting elements) that can be delivered by the k -th supplier of the j -th type for the i -th type of UAV;

$p(j)$ is the number of possible suppliers of the j -th type of weapon.

Then the overall effectiveness of UAVs in a military conflict zone:

$$Q = \sum_{i=1}^N \sum_{j=1}^{n_i} \sum_{k=1}^{p_j} q_{ij} m_{ijk} x_{ijk}, \quad (46)$$

where N is the number of types of UAVs.

To ensure air superiority in the conflict zone, it is necessary to:

$$\max Q, \quad Q = \sum_{i=1}^N \sum_{j=1}^{n_i} \sum_{k=1}^{p_j} q_{ij} m_{ijk} x_{ijk}, \quad (47)$$

taking into account the limitations:

$$T \leq T', \quad T = \sum_{i=1}^N \sum_{j=1}^{n_i} \sum_{k=1}^{p_j} t_{ijk} m_{ijk} x_{ijk}, \quad (48)$$

where T is the time spent on the production and delivery of the UAV and its components to the combat zone;

T' – is the permissible time allocated for the production and delivery of the UAV and its components to the combat zone;

t_{ijk} is the time spent on the production and supply of support for one UAV of the i -th type, j -th type of components by the k -th supplier.

$$W \leq W', \quad W = \sum_{i=1}^N \sum_{j=1}^{n_i} \sum_{k=1}^{p_j} w_{ijk} m_{ijk} x_{ijk}, \quad (49)$$

where W is the cost of production and supply of weapons to the combat zone;

W' – allowable costs;

w_{ijk} – costs associated with the production and supply of one UAV of the i -th type, j -th type of components by the k -th arms supplier.

$$R \leq R', R = \sum_{i=1}^N \sum_{j=1}^{n_i} \sum_{k=1}^{p_j} r_{ijk} m_{ijk} x_{ijk}, \quad (50)$$

where R – risks associated with the production and supply of weapons to the combat zone;

R' – acceptable risks;

r_{ijk} is the risk associated with the production and supply of one UAV of the i -th type, j -th type of components by the k -th arms supplier.

The following condition must be met:

$$\sum_{k=1}^{p_j} x_{ijk} = 1, \quad (51)$$

which means the mandatory selection of a supplier for the i -th type of UAV, j -th type of component.

3.4 Agent-based simulation model for studying the logistics of production and supply of weapons to the combat zone to establish military superiority

To calculate the time required to provide UAV units with the necessary weapons to establish military superiority in the air in the face of threats, a model was created that allows simulating the dynamics of military cargo transportation (requests in the simulation model), taking into account the emergence of threats and the excitation of vulnerabilities, which leads to a halt in the movement of cargo to the combat zone. The agent-based model was created using the Any Logic simulation environment, taking into account the main events that occur during the transportation of military cargo.

The transportation network for the supply of military cargo is represented as a directed graph, in which the vertices are transport hubs and the edges are sections of the transportation highway. Possible transshipment and temporary storage of cargoes are considered in the form of corresponding modeling agents. The main agents in the simulation model for studying the production and logistics of military cargo supply to establish parity of forces in the conflict zone are:

1. Transport network description agent.
2. Request generator agent (military cargo).
3. Threat agent.
4. Vulnerability agent.
5. Military cargo transport route formation agent.
6. Transport hub agent.
7. Transport route section agent.

8. Military cargo temporary storage (warehousing) agent.
9. Transshipment agent.
10. Military cargo time delay agent (due to vulnerability excitation).
11. Risk agent.
12. Combat zone agent (CZA).
13. Simulation control agent.
14. Simulation results agent.

Figure 3 shows a block diagram of the agent-based model.

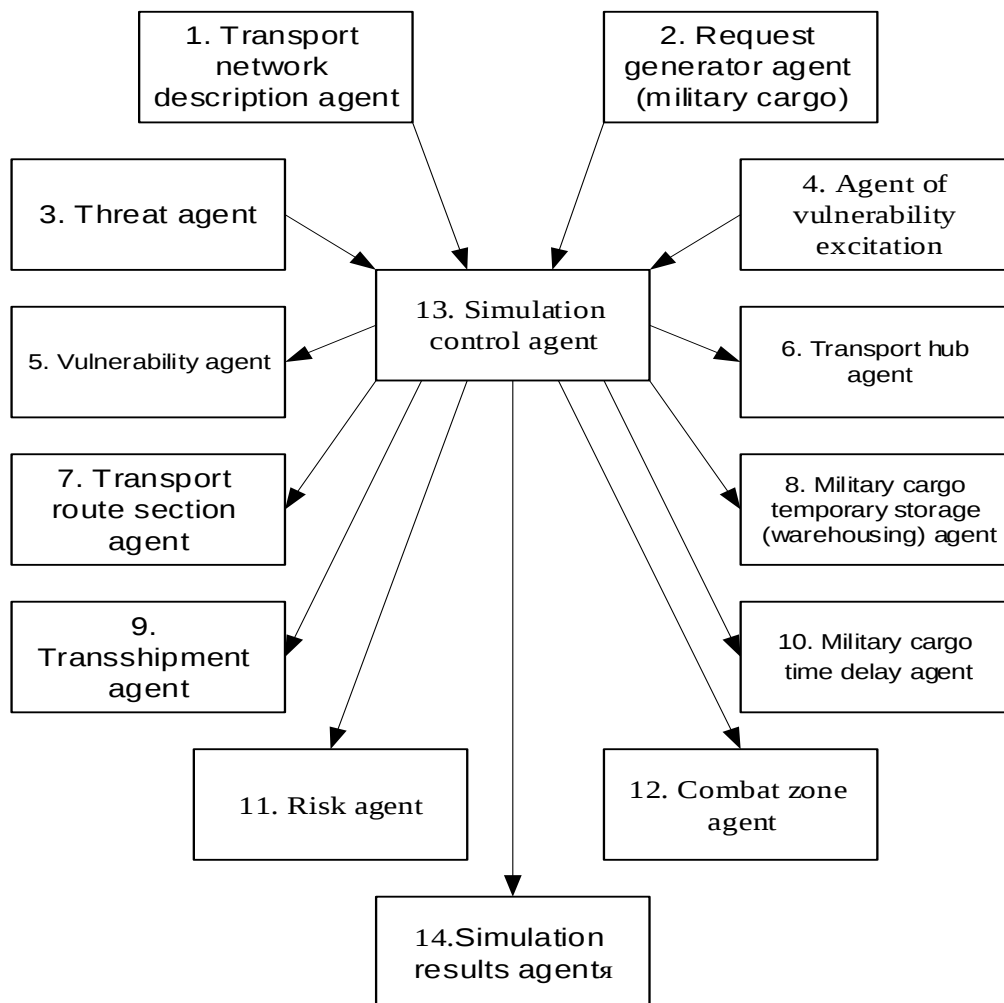


Fig. 3. Block diagram of the agent model

Let us briefly describe the simulation modeling algorithm. With the help of military experts, a transport network (or a fragment of it) is set to be used for the transportation of military cargo to the CZ in the form of sets of transport nodes and sections of a heterogeneous transport network (transport network description agent). Next, the time of the start of the military cargo movement in the transportation highway is set (the request generator agent). Then, with the help of the cargo route

agent, the request (military cargo) is transported through transport nodes and sections of the highway using the transport node agent and the highway section agent. During the movement of a request (military cargo), vulnerabilities (vulnerability agent) may be triggered due to threats (threat agent). This leads to a temporary delay in the application (military cargo) for the period associated with the nature of the threat and the triggered vulnerability. When the request (military cargo) arrives in the combat zone (agent of the combat zone), the simulation stops and the results are generated (agent of the simulation results) in the form of

- the time of receipt of applications (military cargo) in the combat zone;
- time spent on the transportation of military cargo along a given route (without or with vulnerability triggering);
- total time delays of military cargo;
- violation of the terms of delivery of military cargo to the FOB;
- time spent on transshipment;
- time spent on temporary storage (warehousing) of military cargo.

An algorithm has been developed to find a route with the minimum time required to transport military cargo in the CZ. In the case of risk assessments proposed by military experts, an algorithm for the formation of accumulated risk in the logistics of supply in the CZ (risk agent) was developed.

Conclusions

The study is related to the modeling of the integrated supply of military equipment to a military conflict zone in the form of UAVs and supporting components. A representation of military equipment in the form of a component structure is formed. The problems with the integrated supply of various types of UAVs to the military conflict zone are analyzed, which are related to: the variety and multiplicity of types of UAVs and its components, small batches of supplies, complex supply logistics, risks of timely delivery for the formation of military equipment stocks. Using system analysis, a multicomponent structure of military equipment was created, which includes weapons and their main components in the form of a multilayer architecture. The main types of military equipment components are identified: reusable components that ensure unification and reusability in different types of military equipment; new components that need to be developed in projects to create modern equipment; adapted components that can be used in new military equipment by modifying reusable components. The formation of a multilayer architecture of military equipment is carried out taking into account the separate types of components. The creation of stocks of military equipment (various types

of UAV systems and components) in the zone of military conflict is carried out using: army warehouses, manufacturers' warehouses, as well as by producing the required amount of weapons.

An algorithm for the formation of UAV stocks has been developed, taking into account the values of safety stocks necessary to establish military superiority in the air. When creating modified UAV models to be used in a military conflict zone, the component architecture of a new military equipment model is formed using a precedent database that contains possible components for reuse. The search for similar components in the database of precedents required for a new project is carried out with an assessment of the proximity of components by tactical and technical characteristics. Both quantitative and qualitative proximity estimates in the form of linguistic variables are proposed.

An optimization model for selecting logistics channels for the supply of weapons and their components to a military conflict zone, taking into account logistics costs and risks, has been created.

The second part of the study is related to the modeling of threats in the logistics process of forming a military advantage. As a result of the preliminary analysis, the shortcomings of existing methods that consider certain aspects of the process of establishing military parity and advantages that are not linked in a comprehensive solution to the problem are identified. The author analyzes the threat factors associated with the lack of military superiority and parity of forces in the zone of military conflict. Military experts assessed various types of weapons to study their impact on the balance of military forces. The assessment was carried out using qualitative expert assessments in the form of linguistic variables. Given the large number of possible options for choosing suppliers of different types of UAVs, the method of integer (Boolean) programming is used. The criteria for evaluating UAV suppliers are proposed, taking into account the time, costs and risks of production and delivery. The article investigates the sequence: threats – vulnerabilities – losses, which is associated with the emergence of threats in the logistics of production and supply of UAVs to establish military superiority in the air. The author shows how the emergence of a threat affects the creation of vulnerabilities in the logistics of UAV production and supply, which leads to losses in the area of military conflict.

A method for increasing the efficiency of weapons use through innovative types of attack UAVs has been developed, which allows, by creating an asymmetry in military parity of forces, to successfully accomplish the goals of a combat operation in a conflict zone.

An agent-based simulation model was developed to study the logistics process of establishing military superiority in the air. The agent-based simulation model is used to calculate the delivery time and formulate rational routes for the transportation of military cargo, taking into account wartime threats.

The following mathematical and modeling methods were used: system analysis; component design; precedent approach; lexicographic ordering of options; theory of experiments, expert evaluation method, integer (Boolean) programming method; multi-criteria optimization based on linguistic variables, agent-based simulation modeling.

The proposed approach allows, in planning the objectives of a military operation, to formulate requirements for the production and supply of weapons and military equipment, and to plan weapons stocks for effective use (establishing military superiority in the air) in a military conflict zone using various types of UAVs when forming a drone swarm, which contributes to the successful execution of military operations.

This study was funded by the National Research Foundation of Ukraine in the framework of the research project 2023.03/0197 on the topic «Multidisciplinary study of the impact of emergency situations on the infectious diseases spreading to support management decision making in the field of population biosafety».

References

1. Stepaniuk, M. Y., Sinitsyn, I. P., Kotelia, O. V. Problema stvorenniya informatsiynoyi systemy lohistyky v zbroynykh sylakh Ukrainy, shcho vidpovidaye standartam NATO [About applicability of NATO logistics information systems in Ukraine]. Problemy prohramuvannya – Problems in programming, 2018, no. 4, pp. 101–110. DOI: 10.15407/pp2018.04.101
2. Nakonechnyi, O. Analiz umov ta faktoriv, shcho vplyvayut' na efektyvnist' funktsionuvannya systemy lohistyky syl oborony derzhavy [Analysis of conditions and factors influencing the efficiency of the system of logistics of the country defense forces]. Systemy upravlinnya, navihatsiyi ta zv'yazku. Zbirnyk naukovykh prats' – Control, navigation and communication systems, 2019, vol. 3, no. 55, pp. 48–57. DOI: 10.26906/SUNZ.2019.3.048
3. Havrylyuk, I. Yu., Mats'ko, O. Y., Dachkovs'kyi, V. O. Kontseptual'ni osnovy upravlinnya potokamy v systemi lohistychnoho zabezpechennya Zbroynykh Syl Ukrainy [Conceptual basis of flow management in the system of logistic support of the armed forces of Ukraine]. Suchasni informatsiyni tekhnolohiyi u sferi bezpeky ta oborony – Modern Information Technologies in the Sphere of Security and Defence, 2019, vol. 34, no. 1, pp. 37–44. DOI: 10.33099/2311-7249/2019-34-1-37-44

4. Pecina, M., Husak, J. Application of the new NATO logistics system. *Land Forces Academy Review*, 2018, vol. 23, no. 2, pp. 121-127. DOI: 10.2478/raft-2018-0014
5. Fedorovych, O. E., Uruskyi, O. S., Chepkov, I. B., Lukhanin, M. I., Pronchakov, Yu. L., Rybka, K. O., Leshchenko, Yu. O. Modelyuvannya transportnoyi lohistyky viys'kovykh vantazhiv z urakhuvannyam zbytkiv, yaki vynykayut' u zoni boyovykh diy cherez zapiznennya u postachanni [Simulation of transport logistics of military cargo considering the losses occurring in the war zone due to delays in delivery]. *Radioelektronni i komp'uterni sistemi – Radioelectronic and computer systems*, 2022, no. 2, pp. 63–74. DOI: 10.32620/reks.2022.2.05
6. Acero, R., Torralba, M., Pérez-Moya, R., Pozo, J. A. Value stream analysis in military logistics: The improvement in order processing procedure. *Applied Sciences*, 2020, vol. 10, no. 1, article no. 106. DOI: 10.3390/app10010106
7. Fedorovich, O., Pronchakov, Y. Metod formuvannya lohistychnykh transportnykh vzayemodiy dlya novoho portfelyu zamovlen' rozpodilenoho virtual'noho vyrobnytstva [Method to organize logistic transport interactions for the new order portfolio of distributed virtual manufacture]. *Radioelektronni i komp'uterni sistemi – Radioelectronic and computer systems*, 2020, no. 2, pp. 102–108. DOI: 10.32620/reks.2020.2.09
8. Lyu, You., Yuan, Jie-Hong., Sun, Yang., Liu, Jie., Gong, Chun-Ye. Optimization of vehicle routing problem in military logistics on wartime. *Kongzhi yu Juece/Control and Decision*, 2019, vol. 34, iss. 1, pp. 121–128. DOI: 10.13195/j.kzyjc.2017.0983
9. Školník, M. New trends in the management of logistics in the Armed Forces of the Slovak Republic. *Zeszyty Naukowe Akademii Sztuki Wojennej*, 2018, no. 3(112), pp. 53–63. DOI: 10.5604/01.3001.0013.0878
10. Halizahari, M., Daud, Mohamad Faris., Sarkawi, Azizi Ahmad. The Impacts of Transportation System towards the Military Logistics Support in Sabah. *International Journal on Advanced Science, Engineering and Information Technology*, 2022, vol. 12, iss. 3, pp. 1092–1097. DOI: 10.18517/ijaseit.12.3.14516
11. Lai, C.-M., Tseng, M.-L. Designing a reliable hierarchical military logistic network using an improved simplified swarm optimization. *Computers and Industrial Engineering*, 2022, vol. 169, article no. 108153. DOI: 10.1016/j.cie.2022.108153
12. Milewski, R., Smal, T. Decision making scenarios in military transport processes. *Archives of Transport*, 2018, vol. 45, iss. 1, pp. 65–81. DOI: 10.5604/01.3001.0012.0945
13. Nechyporuk, M. V., Fedorovych, O. Ye., Popov, V. V., Romanov, M. S. Modelyuvannya profiliv spetsialistiv dlya planuvannya ta vykonannya proyektiv zi stvorennia innovatsiynykh vyrobiv aerokosmichnoyi tekhniky [Modeling of specialists' profiles for planning and implementation of projects for the creation of innovative products of aerospace techniques] *Radioelektronni i komp'uterni sistemi – Radioelectronic and computer systems*, 2022, no. 1(101), pp. 23–35. DOI: 10.32620/reks.2022.1.02
14. Barbu, M.-L. Theoretical considerations concerning the setting of the capability requirements specific to combat engineers structures supporting management activities from the airfield. *Journal of Defense Resources Management*, 2019, vol. 10, no. 2(19), pp. 188–196.
15. Raskin, L., Sira, O., Katkova, T. Dynamic problem of formation of securities portfolio under uncertainty conditions. *EUREKA: Physics and Engineering*, 2019, no. 6, pp. 73–82. DOI: 10.21303/2461-4262.2019.00985

16. Cappella Zielinski, Rosella., Poast, Paul. Supplying Allies: Political Economy of Coalition Warfare. *Journal of Global Security Studies*, 2021, vol. 6, iss. 1, article no. ogaa006. DOI: 10.1093/jogss/ogaa006
17. Dimitrov, M. S., Irinkov, V. N. State and trends in the development of the logistic system of the Bulgarian Armed Forces. *Obronność-Zeszyty Naukowe Wydziału Zarządzania i Dowodzenia Akademii Sztuki Wojennej*, 2018, no. 3 (27), pp. 35–44.
18. Raskin, L., Sira, O., Sukhomlyn, L., Parfeniuk, Y. Universal method for solving optimization problems under the conditions of uncertainty in the initial data. *Eastern-European Journal of Enterprise Technologies*, 2021, vol. 1, no. 4(109), pp. 46–53. DOI: 10.15587/1729-4061.2021.225515
19. Raskin, L., Sira, O., Palant, O., Vodovozov, Y. E. Development of a model of the service system of batch arrivals in the passengers flow of public transport. *Eastern-European Journal of Enterprise Technologies*, 2019, vol. 5, no. 3(101), pp. 51–56. DOI: 10.15587/1729-4061.2019.180562
20. Fedorovych, O. Ye., Rubín, E. Ye., Yeremenko, N. V. Issledovaniye logistiki snabzheniya i sbyta v raznorodnoy transportnoy infrastrukture gruzoperevozok [Study of supply and marketing logistics in a heterogeneous transport infrastructure of cargo transportation]. Kharkiv, Nats. aerokosm. un-t «Khark. aviats. in-t», 2016. 198 p.
21. Fedorovych, O., Polishchuk, Ye., Chmykhun, Ye., Solovyov, V. Modelyuvannya krytychnykh vrazlyvostey u lohistytsi postachannya ozbroyennya ta viys'kovoyi tekhniky v umovakh voyennykh zahroz [Simulation of critical vulnerabilities in the logistics of supply arms and military equipment under conditions of military threat]. *Aviacijno-kosmicna tehnika i tehnologia – Aerospace technic and technology*, 2022, no. 6 (184), pp. 40–49. DOI: 10.32620/aktt.2022.6.05

**DIGITIZATION OF IT TEAM MANAGEMENT PROCESSES
IN PROJECT MANAGEMENT:
THE ROLE OF INFORMATION SYSTEMS
AND ARTIFICIAL INTELLIGENCE**

Kovalchuk O., Kobylkin D.

The monograph is devoted to the study of digitalization processes in IT team management within project and program management, with a focus on the role of information systems and artificial intelligence. The work examines the theoretical and methodological foundations of digitalization, the evolution of management approaches from classical models to Agile, DevOps, and scaled frameworks such as SAFe. Special attention is given to the analysis of project and HR management information systems, as well as the integration of AI in planning, productivity forecasting, resource optimization, and talent development processes. The authors present a team life cycle model, highlighting key characteristics, challenges, and managerial tasks at various stages of team formation, growth, stabilization, and renewal. Global and Ukrainian practical cases of digitalization and the implementation of innovative technologies are presented, demonstrating opportunities to enhance efficiency, transparency, and adaptability of teams in dynamic IT environments. The monograph may be useful for researchers, educators, IT team leaders, HR specialists, and managers aiming to apply modern digital approaches to improve organizational performance.

Current trends in society and the economy dictate the need to transform project management systems through the active implementation of digital technologies. Human resource management (HR) is becoming particularly important, as human capital remains a key factor in the success of any project. The use of information systems and artificial intelligence (AI) tools opens up new opportunities for optimizing the processes of selection, adaptation, motivation, and development of personnel within the implementation of projects and programs.

Adaptive Agile methodologies for managing IT projects and programs stimulate innovative team development and the search for new ideas. Masaaki Imai, a well-known theorist and consultant in the field of quality management, developed the KAIZEN concept. After World War II, management practices were introduced in organizations, including Toyota, and contributed to Japan's economic miracle. This management system is aimed at the continuous improvement of processes in the organization and its functional aspects. This method belongs to the group of total quality management (TQM) systems. It can positively influence the processes of building an HR management system, as its principles encourage each team member to contribute to the achievement of project goals. A systematic approach

to the implementation of minor stages has a positive impact on the overall success of an IT project. KAIZEN practices also help to identify and eliminate project bottlenecks and losses, namely financial, time resources, and team member efforts.

KAIZEN is combined with Agile methodologies and the Scrum method. This allows for improved efficiency in team member communication. Daily sprints, retrospectives, and meetings facilitate problem solving, progress monitoring, rationalization proposals, and staff involvement in the implementation of IT projects and programs.

According to Isaac Adizes' theory, the life cycle of a team is characterized by certain stages that reflect the development of the team from the moment of its formation to the possible termination of its activities or restructuring. The first stage is the birth of the team, when its foundations are laid and the initial idea is formed. During this period, goals and initial resources are defined, and a leader and shared values are sought to determine the future direction of the team's work. The main tasks at this stage are to create the initial structure of the team and distribute key functions among its members. At the same time, the team faces a number of challenges, including uncertainty, lack of resources, and the need to establish clear roles.

The second stage is the accelerated growth or youth of the team, when the team rapidly increases in size and volume of work performed. At this stage, there is a need to professionalize participants, standardize processes, and implement a management system to avoid chaos. The main tasks are to build an effective organizational structure, distribute responsibilities, and adapt to new challenges that arise in the growth process. The youth stage is critical for consolidating the foundation of team effectiveness, as it is during this period that the standards of cooperation and interaction are established that will influence the further development of the team.

The third stage is stabilization or maturity, during which the team achieves stability and efficiency, and its activities become recognizable and predictable. The main tasks of this period are to maintain stability and efficiency while simultaneously searching for new avenues of development. The team must avoid the risk of complacency and loss of flexibility, which can lead to stagnation. Innovation and adaptation to changes in the external environment become key factors in maintaining the viability of the team at this stage.

The fourth stage is aging or renewal of the team, when productivity may decline and signs of internal degradation may appear. At this stage, there is a need for fundamental changes and a review of the team's goals and values. The main tasks are to update the structure, define a new mission, or, if reform is not possible,

prepare for the end of the team's activities. The aging stage demonstrates the importance of strategic management and the organization's ability to transform its teams to ensure long-term effectiveness.

Thus, the team life cycle model reflects the dynamics of its development, outlines the key characteristics of each stage, potential challenges, and necessary management decisions. Understanding these patterns allows managers to effectively plan resources, develop team members' competencies, and ensure long-term team effectiveness, which is especially relevant for IT teams working in project and program management, where the speed of change and complexity of tasks are increasing significantly. The team life cycle according to Adizes is a dynamic process where each team must be aware of its current position and actively work to overcome the challenges of each stage to achieve further success.

By applying Kaizen principles, HR managers can transform themselves from «administrators» to «facilitators» of change, helping teams to self-organize and continuously improve, thereby moving between the stages of the team-building life cycle. This not only increases efficiency but also strengthens corporate culture.

In today's digital transformation environment, organizations are faced with the need to simultaneously manage multiple teams working on complex information products and services. The growth in the scale of digital initiatives, the increase in the number of project participants, and the complexity of inter-team interactions require organizations to adopt new approaches to planning, coordination, and control. Traditional management methods no longer provide the necessary level of flexibility and speed of response to changes in the market environment. That is why digitalization, supported by the introduction of information systems and artificial intelligence tools, is becoming a determining factor in improving the efficiency of modern management processes.

Traditional Agile approaches, such as Scrum or Kanban, have proven themselves well within small teams, but their application becomes limited when scaled to the program or portfolio level. To address this issue, the Scaled Agile Framework (SAFe) was developed, which is a methodological approach to the coordinated management of a large number of teams within a single product vision.

Digitalization in team management is interpreted as the process of integrating digital technologies into all stages of a project or program lifecycle, which allows for the optimization of resource allocation, ensures process transparency, and improves decision-making quality. In the context of IT teams, the evolution of management methods is of particular importance: from classic hierarchical models to modern flexible methodologies such as Agile and DevOps. Agile methodology

has provided a new level of flexibility in managing small teams, but with the growth of organizations, there has been a need for methods capable of integrating the work of dozens or even hundreds of teams within a single strategic framework. The response to this challenge was the formation of scalable models, among which the Scaled Agile Framework (SAFe) occupies an important place.

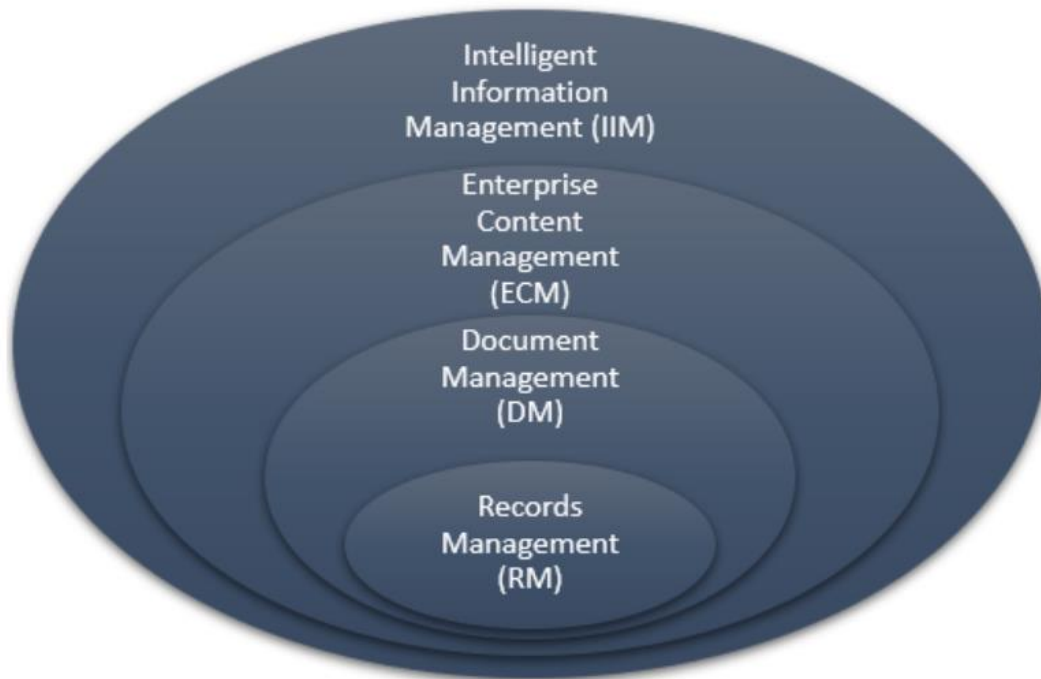


Fig. 1. Intelligent Information Management

SAFe combines Lean, Agile, and DevOps principles, allowing you to strike a balance between task flexibility and strategic organizational goals. The framework is based on the Agile Release Train (ART) concept – a long-term cross-functional structure that brings together several teams to work together on a product or program. ART ensures synchronization of iterations, integration of results, and increased transparency in the development process.

The SAFe structure provides for four levels of management:

- Team level – a separate team working according to Scrum or Kanban, ensuring the completion of tasks within sprints.
- Program level – coordination of several teams within the ART, allowing them to coordinate their work and jointly achieve intermediate results.
- Large Solution level – management of large systems that require the interaction of several programs. This level provides scalability and integration at the level of architecture and technological solutions.

- Portfolio level – a strategic level that combines projects and programs with common business goals, sets priorities, and allocates resources in accordance with the company’s long-term strategy.

An important advantage of SAFe is its ability to ensure consistency between business and technical teams, which is particularly relevant for the IT sector. Thanks to SAFe, organizations can combine strategic planning with operational flexibility, minimize risks, and respond more quickly to changes in the market environment.

The digitization of HR processes involves the comprehensive use of information systems to collect, process, and analyze employee data. This reduces routine tasks and allows you to focus on strategic aspects of management. The use of specialized HRM systems is integrated with project management systems, allowing you to form optimal teams according to competencies and tasks.

The use of information systems within SAFe plays a key role. Project and program management systems (Jira Align, Rally Software, Microsoft Azure DevOps) allow you to visualize planning, track dependencies between teams, and automate reporting processes. In addition, the use of artificial intelligence tools in SAFe helps to improve the accuracy of risk forecasting, optimize resource planning, and ensure a data-driven approach to management.

Thus, SAFe is not only a methodology for scaled Agile, but also a comprehensive approach to managing large project portfolios, combining organizational strategy with flexible teamwork practices. Its implementation in IT team management ensures synchronization, transparency, and speed of management decision-making, which is critically important in today’s digital environment.

Table 1

Global cases and the Ukrainian context: digital transformations

Category	Case / Example	Focused on
World SAFe	Toptal: Telecom Serbia	Training, team coordination, Agile Release Train
World SAFe	Toptal: financial company	WSJF, prioritization, synchronization implementation
Ukrainian case studies	Air Alert, Diia, Prozorro.Sale	Agile MVP, digital services, e-government, transparency
Ukrainian case studies	AI-проекти (Noty.ai, Signy)	Task automation, AI assistants, SaaS

SAFe combines Lean, Agile, and DevOps principles, providing organizations with a structured model for managing large programs and portfolios. Its distinctive

feature is a multi-level management structure that covers the team, program, large solution, and portfolio levels. At the team level, SAFe involves working in Scrum or Kanban format; at the program level, it involves coordinating multiple teams through Agile Release Train; at the large solution level, it focuses on integrating complex systems; and at the portfolio level, it combines projects with long-term business goals. Thanks to this, SAFe ensures synchronization between the organization's strategic priorities and the operational activities of teams, which is especially important for the IT sector.

Information systems are an important element of digitalization management. They serve as a basic tool for organizing teamwork, coordinating tasks, and managing resources. Among the most common systems are Jira, Trello, Asana, Monday.com, Microsoft Project, and ClickUp, which provide task management, progress tracking, integration with third-party tools, and automation of some processes. More complex solutions, such as Jira Align or Rally Software, allow organizations to work according to SAFe principles and coordinate a large number of teams within a single program. Information systems digitize key management functions – planning, monitoring, control, and reporting – enabling managers to respond more quickly and accurately to changes in the external and internal environment. To eliminate IT project «losses», successful organizations use digital services to automate HR processes: recruiting, onboarding, talent development, and retention. Artificial intelligence is beginning to play a special role in modern team management. Artificial intelligence in human resource management is used for:

- analyzing large amounts of data about candidates and employees;
- predicting productivity and staff turnover risks;
- creating individual staff development trajectories;
- supporting management decision-making.

In project management, this allows for more accurate planning, shorter task completion times, and reduced risks associated with inefficient use of human capital. Its integration into digital management systems allows for the automation of a significant number of routine tasks: resource allocation, priority setting, task completion time forecasting, risk identification, and team performance analysis. Machine learning algorithms provide deeper analysis of large data sets and offer recommendations for decision-making based on predictive analytics. In the context of IT team management, this means the ability to proactively identify bottlenecks, timely adjust team workloads, and optimize communications. At the same time, the use of AI requires attention to issues of ethics, data protection, and preventing excessive reliance on algorithms in the strategic decision-making process.

There are numerous examples of successful application of SAFe and information systems in combination with AI in global practice. For example, international corporations such as Intel, AstraZeneca, Royal Philips, and Capital One demonstrate the effectiveness of scalable Agile approaches in coordinating the work of thousands of employees. Individual studies show that thanks to SAFe, companies reduce time to market, improve alignment between business strategies and technical solutions, and increase transparency for customers. In the financial and telecommunications sectors, SAFe has solved the problem of inter-team fragmentation by introducing mechanisms for synchronizing and prioritizing tasks using the WSJF model.

Ukraine also has examples of large-scale digitalization of government, which, although not always directly related to SAFe, demonstrate the effectiveness of Agile and AI in large projects. In particular, the Diya platform has become a symbol of digital transformation in public administration and was implemented according to the principles of MVP and rapid scaling. The Air Alert system, developed during military operations, demonstrated the capabilities of rapid iterative development and integration of user needs in the shortest possible time. Prozorro.Sale showed how digital tools can ensure transparency and efficiency in the management of state assets. These examples show that the digitization of management can be successfully applied not only in the commercial sector but also in public administration, opening up new opportunities for the use of information systems and artificial intelligence.

In today's digital world, information systems and artificial intelligence are becoming critical tools for human resource management, especially in the IT sector, where the speed of change and complexity of tasks require teams to be highly adaptable. Modern HR systems allow for the automation of a wide range of processes that were previously performed manually, including recruitment, onboarding of new employees, career planning, performance evaluation, and knowledge management. Integrating these systems into project and program management provides more accurate control over resources, effective task distribution, and optimization of teamwork. Modern HR information systems such as SAP SuccessFactors, Workday, Oracle HCM Cloud, BambooHR, and Personio allow you to collect and analyze large amounts of employee data. These platforms provide centralized personnel management, KPI tracking, performance monitoring, and competency development planning in line with the organization's strategy. Thanks to such systems, managers can assess team performance in real time, identify bottlenecks in task execution, and

make informed management decisions, which significantly increases the team's adaptability at all stages of its life cycle.

Artificial intelligence in HR opens up new opportunities for automating decision-making processes. Machine learning algorithms allow you to predict the effectiveness of candidates at the recruitment stage by analyzing both professional and behavioral characteristics. AI assistants help automate the onboarding process for new employees by selecting individual training plans and assessing progress in mastering the necessary skills. In addition, intelligent systems analyze team dynamics, identify potential conflicts, and provide recommendations for optimal task distribution, which helps increase productivity and employee satisfaction. In project and program management, the integration of HR information systems and AI contributes to more effective human resource management in large-scale IT projects. For example, automated systems can track team members' workloads, predict task completion times based on historical data, and suggest optimal resource allocation scenarios. This is especially important during periods of rapid growth and team stabilization, when the number of tasks and participants increases and the risk of overload and loss of efficiency becomes critical.

The use of AI in HR also promotes the development of a personalized approach to talent management. Intelligent platforms are capable of creating individual development plans for employees, identifying training and upskilling needs, and assessing the alignment of current competencies with the organization's strategic goals. This ensures the continuous development of the team, maintains its motivation, and prevents the risk of stagnation, which is characteristic of the maturity stages of the team life cycle. Special attention should be paid to data analytics in HR, which allows managers to obtain data-driven insights into team performance. Based on large amounts of information about productivity, participant interaction, and project results, AI systems can generate forecasts of potential productivity declines, suggest measures to optimize processes, or even identify hidden talents among employees. This approach not only helps maintain high team performance but also contributes to the strategic development of the organization as a whole.

Thus, the introduction of information systems and artificial intelligence in HR allows for the creation of a new generation of human resource management systems, where the automation of routine processes is combined with deep analytics and the individualization of management decisions. For IT teams working in project and program management, this means increased adaptability, efficiency, and innovation at all stages of the life cycle, from the team's inception to its stabilization and renewal. The integration of HR information systems and AI also opens up

prospects for further scientific research, in particular regarding the evaluation of the effectiveness of algorithms in predicting team performance, studying the impact of automated decisions on employee motivation, and developing ethical standards for the use of artificial intelligence in human resource management. This allows for the combination of strategic planning and digital technologies, creating a comprehensive approach to managing IT teams in today's environment.

Thus, the digitization of IT team management processes in project and program management is a complex phenomenon that encompasses methodological, technological, and organizational aspects. It is based on a combination of modern approaches, such as Agile, DevOps, and SAFe, with the use of powerful information systems and artificial intelligence tools. This allows organizations to increase their adaptability to change, reduce project implementation time, improve communication between teams, and ensure a higher level of transparency in management processes. Prospects for further research lie in studying the practical aspects of integrating AI into program management systems, developing ethical standards for the use of algorithms in management, and adapting scalable approaches to the specifics of the Ukrainian IT sector. Therefore, the digitization of HR processes and the use of artificial intelligence are becoming key factors in improving project management efficiency. The integration of information systems into project management contributes to the creation of flexible teams, the optimization of management decisions, and the growth of organizations' competitiveness.

References

1. IT Management: Textbook: in 2 parts. Part 1 / O.V. Gorpinchenko, O.V. Zayarnuk, I.M. Sochynska-Sybirtseva [et al.]. Kropyvnytskyi: CNTU, 2024. – 218 p.
2. Balanovska T. I., Mikhailichenko M. V., Troyan A. V. Modern technologies of personnel management: textbook. Kyiv: FOP Yamchynskyi O. V., 2020. 466 p.
3. N. Kovalchuk, O. Zachko, O. Kovalchuk and D. Kobylkin, «Project Management of the Information System for the Selection of Project Teams», 2023 IEEE 12th International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS), Dortmund, Germany, 2023, pp. 1054–1057, DOI: 10.1109/IDAACS58523.2023.10348656
4. Agres, O., Sodoma, R., Ilchyshyn, I., Kovalchuk, O., & Shmatkovska, T. (2025). Regional development project management: financial aspect. *Technology Audit and Production Reserves*, 3(4(83)), 87–92. <https://doi.org/10.15587/2706-5448.2025.330027>
5. Kovalchuk Oleh, Kobylkin Dmytro and Zachko Oleh: Graphodynamic modeling for a multi-agent support system for personnel decision-making in the field of human safety. *Proceedings of the 4th International Workshop IT Project Management (ITPM 2023)*. Warsaw 2023. P. 149–159.

**MANAGEMENT OF MONITORING AND CONTROL PROCESSES
IN INNOVATIVE PROJECTS OF UAV PARAMETERIZATION
AND SWARM FLIGHT: MODELS, STRATEGIES, ADAPTIVE SOLUTIONS**

Kritskiy D., Malyeyeva O., Popov A., Gubka O., Kosenko N.

Within the framework of managing innovative projects for the development of autonomous aviation systems, a comprehensive approach to the parameterization of unmanned aerial vehicles (UAVs) and the creation of group flight control models in swarm structures has been investigated. The aim of the work is to ensure the effective and safe operation of UAVs in a swarm by calculating their parameters and controlling group flight in a changing air environment based on a mathematical and knowledge-oriented model. The following tasks are addressed: analysis of autonomy levels, identification of critical parameters of devices, modeling of airspace with conflict assessment, and construction of airspace ontology. Analyze the levels of autonomy and control in a UAV swarm; determine the key parameters of UAVs that affect their interaction; build a mathematical model of airspace taking into account potential conflicts; develop an ontological model of airspace zones taking into account accessibility categories and restrictions. Results: The main focus is on the formation of design solutions for autonomous interaction between devices, multi-level control in a swarm, and air environment modeling. Mathematical models and a knowledge-based database have been developed to support innovative solutions for the safe operation of UAVs in dynamic environments. The results can be integrated into innovation management programs in the field of defense technologies, aviation design, and intelligent transport systems..

Introduction

Unmanned aerial vehicles (UAVs) have been in increasing demand in recent years for civil, commercial, and military applications. In military applications in particular, drones are becoming an integral part of the future battlefield [1]. Due to the limitations of the payload and energy of a single aircraft, it has been proposed to deploy swarms of drones for a variety of applications with higher mission requirements, which is changing their future application [2]. Drone swarms perform missions jointly and have the potential for rapid deployment and coverage that goes beyond the individual capabilities of single units. Currently, the number of simultaneous drone missions is estimated to be in the tens or hundreds.

The large-scale use of drones is a new innovative element of modern hybrid warfare. The formation of strike drones into swarms, the planning of their missions in waves, and the provision of adaptive control and coordination allow for the mass destruction of relevant enemy targets and contribute to the successful execution of operational and tactical actions by military leadership. The urgency of these issues

is caused by the current military and political realities in Ukraine and corresponds to urgent state priorities and needs.

In the current conditions of rapid development of autonomous systems, there is a growing need for effective management of group flights of unmanned aerial vehicles [3]. Such swarm structures allow for monitoring, search, mapping, and reconnaissance tasks to be performed with minimal operator intervention. However, the effective organization of a UAV swarm requires clear structuring of the aircraft, determination of their technical and functional parameters, and the construction of interaction models in a complex, changing air environment.

New adaptive and intelligent methods for planning and controlling combat drones are one of the most important components of modern hybrid warfare, allowing for increased military effectiveness in conditions of limited time and increased risk.

The intensive development of UAVs and the expansion of their areas of application, particularly in the military sector, requires the introduction of innovative solutions in the field of complex technological systems management. One promising area is the use of UAV swarm systems, where a group of autonomous devices performs joint tasks based on distributed interaction. Effective management of such systems requires coordinated models for parameterization, data exchange, decision-making under uncertainty, and the ability to adapt to a dynamic environment.

One of the key challenges is the formalization of models of autonomy, control levels, and interaction scenarios for UAVs within a swarm. At the same time, it is necessary to model the airspace, taking into account potential conflicts and restricted access zones. To support adaptive planning and decision-making in conditions of many variables, an approach is used to construct an ontological model of flight zone categories.

In view of this, there is a need to develop a comprehensive approach to managing innovative projects related to the development of UAV swarm systems. Such projects combine elements of system analysis, mathematical modeling, engineering design, and knowledge-based technologies. The creation of an effective group flight model that takes into account changes in the air environment, potential conflicts, and restrictions on airspace availability allows for high mission efficiency.

The development of mathematical and ontological models within such projects allows for the optimization of planning, monitoring, and management processes for innovative developments in the field of autonomous aviation systems. At the same time, ensuring the safety and reliability of swarm solutions is a critically important step on the path to the practical implementation of such technologies. That is why research aimed at parameterizing UAVs and

formalizing group flight control models is an important component of strategic management of innovative development in high-tech industries.

1. Analysis of autonomy, control levels, and interaction of UAVs in a swarm

The autonomy of UAVs is a measure of independence from external influences and the level of self-organization [4, 5] (Table 1). The first three levels (from 0 to 2) correspond to UAVs that require constant human operator involvement. They use systems that facilitate control, but constant analysis of the external situation is necessary. Levels 3–5 of autonomy correspond to UAVs that are already capable of moving independently and analyzing data obtained from the environment. Control of level 3 and 4 drones can be transferred to a human operator if necessary, while level 5 corresponds to UAVs that operate exclusively in autonomous mode.

Table 1

Levels of UAV autonomy

Level of autonomy	0	1	2	3	4	5
Degree of automation	None	Low	Part	Conditional	High	Full
Avoiding obstacles	None	Recognize and warn		Recognize and avoid	Recognize and orient	

Let's consider the levels of UAV autonomy and how the implementation of autonomous UAV control looks in practice:

Level 0. No automation. The pilot has complete control over every movement. The platforms are constantly monitored 100% manually.

At zero autonomy, no function is performed by the drone independently. All tasks, including maintaining altitude and the required flight speed, are the sole responsibility of the operator. Such UAVs are very rare and are not intended for swarm use. They are mainly used for racing or are the simplest UAVs loaded with cameras for photo or video recording. Only an experienced operator can successfully control such a drone.

Practical application: UAVs for model aircraft.

Level 1. Pilot assistance. The pilot remains in control of the overall operation and safety of the UAV. However, the UAV may take over at least one vital function for a certain period of time.

The first level of autonomy involves the UAV performing only basic functions, such as automatically maintaining altitude or flight speed. Satellite or inertial

navigation aids may be integrated into the control system, with data from these being analyzed by the operator or a ground station. All key flight phases, from takeoff to landing, remain under human control. Drones at this level may be equipped with sensors and cameras, but are not capable of processing the information they receive independently. They support GNSS for flight stabilization, while all inputs in terms of course, altitude, and speed are made manually. At this level, collision detection and avoidance functions are available to warn the pilot of nearby UAVs and other obstacles, but some functions depend on manual control by the pilot. Due to their limited capabilities, swarm systems cannot be created with these UAVs.

Applications: Mainly used for entertainment purposes, surveillance, photography, and videography. Or used to obtain simple recreational photos and videos.

Military applications: localization and detection, photography and videography, protection and monitoring.

Level 2. Partial automation. The pilot is still responsible for safe operation and must be ready to take control of the UAV if something happens. However, under certain conditions, the UAV can take control of the course, altitude, and speed. The pilot monitors the airspace, flight conditions, and responds to emergencies. Currently, most manufacturers are developing UAVs at this level, where the platform can assist with navigation functions and allow the pilot to relinquish some of their tasks.

At the second level of autonomy, the UAV is capable of performing certain functions without operator intervention, such as stabilizing its position, following a predetermined route, or keeping an object in the frame. However, overall flight control and decision-making remain with the human operator. Data from navigation systems (GPS, inertial sensors) can already be processed partially automatically, but tasks such as obstacle avoidance still require operator intervention.

Practical application: a pre-programmed flight path is loaded onto the autopilot and the UAV begins to perform tasks along these trajectories after takeoff. This is now common practice in mapping. Some UAVs with this level of autonomy have a built-in automatic takeoff and landing function, which makes the UAV easier to control but does not make it more autonomous.

Military applications: mapping, surveying. This level allows UAVs to be used for simple reconnaissance, monitoring, or delivery missions, but such drones are extremely limited in swarm operations.

Level 3. Conditional automation. As in level 2, the UAV can fly on its own, but the human pilot must still remain alert and be ready to take control at any time. The UAV control system transmits data to the pilot about the need for intervention,

so it is a backup system. This level means that the UAV can perform all functions «under certain conditions».

At the third level of autonomy, UAVs gain obstacle recognition, which complements the capabilities of the previous level. This allows the drone to analyze its environment, but it cannot yet make decisions. When obstacles are detected, the device stops and notifies the operator of the need for intervention.

In swarm interaction, two approaches are provided for UAVs of this level. The first is control via a leader drone. In such a system, a person or ground station controls one main drone, and all others follow it, having a direct connection with the leader. At the same time, the distance from it should not exceed the stable communication range. To avoid loss of control, when communication is lost, a new leader is automatically selected from among the swarm members.

The second option is centralized control from a command center. In this case, each UAV sends all data (coordinates, speed, direction of movement, acceleration, sensor information, etc.) to the control center. The center analyzes the situation and sends commands to the drones. With this architecture, there is no interaction between the UAVs, and all decisions are made externally, not on board the aircraft.

Practical application: The UAV flies along a programmed flight path until its onboard sensors detect an obstacle. Upon detection, the UAV stops and sends an alarm signal about the object in its immediate vicinity. The pilot then manually corrects the direction and altitude before the UAV continues flying along its programmed path.

Military applications: mapping, monitoring, and cargo delivery. Swarm systems with third-level autonomy drones have sufficient functionality for mapping and security monitoring tasks.

Level 3+. Automated UAV deployment. Another way to measure UAV autonomy is to consider the operating environment. Some manufacturers have made progress in automating UAV deployment. This means that there is no human in the loop to control the flight.

Practical application: the idea is to have a UAV installed on site that frequently performs the same task (e.g., surveying a specific area twice a day). The pre-programmed flight path remains unchanged, the UAV is equipped with automated take-off and landing capabilities, and the landing site allows the batteries to be recharged and the UAV to be automatically deployed according to a set schedule.

Military applications: mapping, surveying, and protection and security.

Level 4. High automation. The UAV can be controlled by a human, but this is not always necessary. It can fly for long periods of time under certain circumstances. UAVs are expected to have backup systems, so if one system fails, it will still function. The behavior of a UAV depends on fixed built-in functionality or a fixed set of rules that dictate the behavior of the system.

The fourth level of autonomy is currently the most advanced among existing UAVs. Drones of this level are capable of independently processing information received from installed equipment and making decisions based on it. For example, when obstacles are detected, they can change their route without external commands and successfully complete their tasks autonomously. However, in non-standard or complex situations, drones may not achieve the planned result. Such UAVs can be controlled either manually or completely autonomously.

Level 4 drone swarms operate independently of ground control. Their autonomy allows them to exchange data with each other, including only with the nearest devices. This optimizes energy consumption and reduces the technical complexity requirements for each UAV. Such a swarm can include drones of different models and with different equipment. Tasks are distributed according to the capabilities of each drone and its specialization.

Bio-inspired algorithms based on the behavior of living organisms are often used to create mechanisms for interaction between drones of this level. These can be algorithms that mimic the actions of ants, bees, bats, or even simple biological particles. Despite different approaches to solving problems, they share a common idea: each drone independently searches for a way to achieve its goal using information received from other members of the swarm. Over time, as the algorithm works, the drones concentrate on the most promising objects for achieving the goal, such as the closest, most valuable, or most important according to specified criteria.

Practical application: A UAV equipped with a specific set of sensors monitors an object. In doing so, the UAV senses obstacles in its flight path and actively avoids contact by changing direction and altitude. It finds a way to complete the task while automatically changing the way it is performed.

Military applications: photography and video recording.

Level 5. Full automation. The UAV controls itself under all circumstances without waiting for human intervention. This includes full automation of all flight tasks under all conditions.

Such UAVs must use AI tools to plan their flights, in other words, autonomous learning systems with the ability to modify routine behavior functions. It is also important to remember that the fifth level of autonomy is a prerequisite

for achieving two key goals in the unmanned industry: air mobility and large UAVs for cargo delivery.

The fifth level of UAV autonomy remains hypothetical and has not yet been implemented in practice. It is assumed that unmanned aircraft of this level will be able to perform assigned tasks completely autonomously in any environment, regardless of external factors. Their operation will be based on highly developed artificial intelligence technologies that will allow drones to independently analyze situations, make decisions, and adapt to changes.

The possibility of operator intervention will be implemented only as an additional backup mechanism, not as the primary form of control.

The communication architecture in swarms of drones at this level will likely be similar to that used in fourth-level UAVs: data exchange between devices will occur in a decentralized manner, with a preference for local interaction between the closest members of the swarm.

The first three levels (0 to 2) correspond to UAVs that require constant human operator involvement. They use systems that facilitate control, but constant analysis of the external situation is necessary. Autonomy levels 3–5 correspond to UAVs that are already capable of moving independently and analyzing data obtained from the environment. Control of level 3 and 4 drones can be transferred to a human operator if necessary, while level 5 corresponds to UAVs that operate exclusively in autonomous mode.

In accordance with the classification of UAV autonomy levels, let us consider *the levels of UAV control* (Fig. 1).

At the first level, interaction with the environment occurs through the control and modification of the operation of the executive mechanisms.

At the second level, the ability to control the rotation speed of the propellers is added to the control of the executive mechanisms.

Today, the simplest flight controllers include a microcontroller with a minimum set of sensors for assessing position, orientation, and altitude (accelerometers, gyroscopes, barometers), which allow you to assess the position and automatically influence the propeller rotation speeds, for example, to stabilize the position in the air and maintain altitude.

The third level controls the orientation of the UAV, whose controller is equipped with a GPS sensor that allows it to orient itself in space.

At the fourth level, linear speed control is automated. In addition to the angular speed and orientation controllers, the autopilot has a level for controlling linear speeds and position.

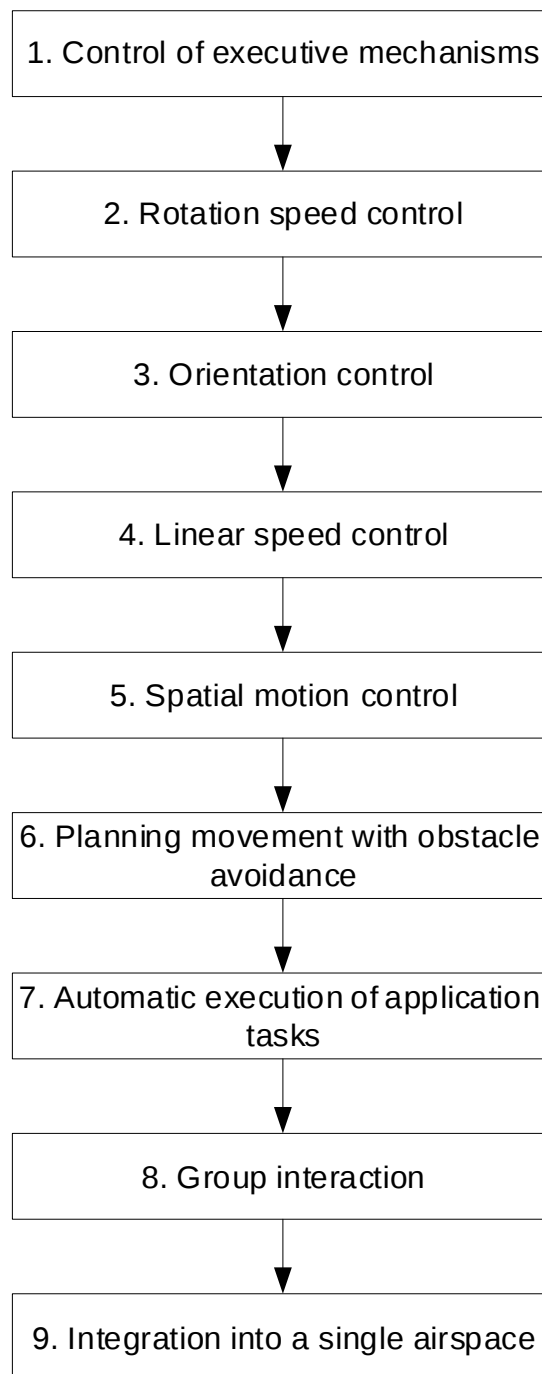


Fig. 1. Levels of UAV motion control

At the fifth level, the movement of the UAV in space is controlled by software. The principle of position control in autopilot mode is as follows: the operator sets the coordinates of a virtual point in space, the autopilot calculates the direction relative to the current coordinates of the UAV and forms a velocity vector, which is sent to the speed controller. The speed controller forms the desired angle of the UAV, which is sent to the controller responsible for the angle. A properly configured control system allows stable flight in windy weather, compensating for external influences.

At the sixth level, intelligent planning systems are integrated into the UAV control system. The UAV must be equipped with algorithms for obstacle avoidance or independent reconnaissance of complex terrain. For this purpose, an on-board computer with sensors is additionally used: cameras, lidars, radars, and other additional devices.

At the seventh level, automatic execution of applied tasks appears. The tasks of the UAV include independent route planning: for example, following any object or automatically exploring and mapping the space.

The eighth level of control is provided by group interaction between UAVs.

The ninth level involves built-in air traffic control systems, as described in [6].

Here is a classification of UAV groups based on **the types of interaction** between devices [7, 8]:

Level 1. Direct physical interaction. In this case, UAVs are connected directly to each other and their movement is limited by forces that depend on the movement of other UAVs. An example is the lifting and transport of cargo by several UAVs. From the point of view of route planning and collision avoidance, all agents can be considered as a single entity. Since the number of devices in this case is usually small, both centralized and decentralized control systems can be used.

Level 2. Formations. The devices are not directly connected to each other, and their relative positions are strictly defined to maintain the formation. Path planning can be considered for a formation of several UAVs as a single entity. Collision avoidance between devices can be built into the formation control algorithms. Scalability becomes important for this approach, and therefore decentralized control algorithms are preferred.

Level 3. Swarms of devices – a team of multiple UAVs, whose interaction algorithms ensure collective behavior. The resulting movement of the group is not necessarily a formation. Scalability is one of the most important issues in this approach, so the use of decentralized algorithms is a prerequisite.

Level 4. Cooperation. UAVs from a group plan their movement according to individual tasks that must be distributed to perform a higher-order mission in the control system hierarchy [9]. These trajectories are usually geometrically interrelated, as in the case of formations. Therefore, issues such as task allocation, high-level planning, plan breakdown into subtasks, and conflict resolution must be resolved before the higher-order mission can be executed. Accordingly, both centralized and decentralized control system architectures must be used.

Figure 2 shows a schematic interaction between the UAVs described above.

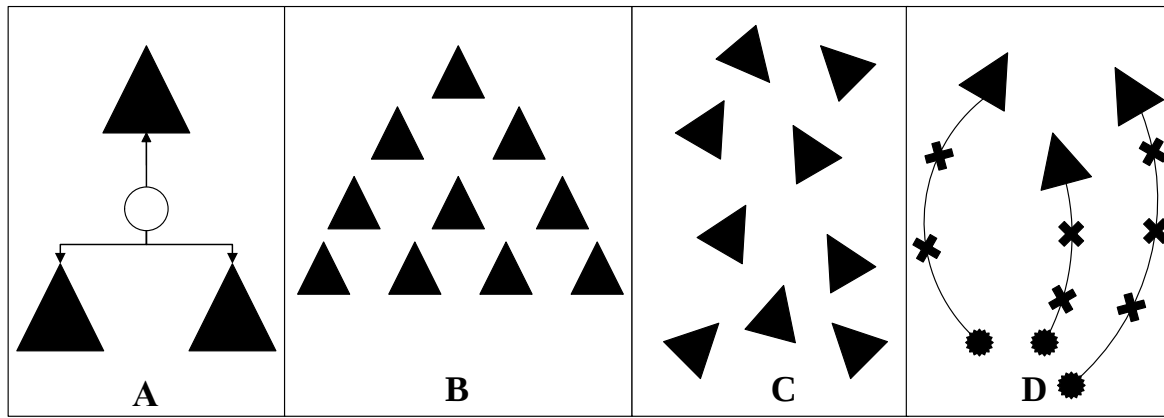


Fig. 2. Types of interactions between UAVs:

A – direct physical interaction, B – formation, C – swarm, D – cooperation

Three possible control strategies can be identified:

- centralized – remote control with a dedicated base station, the swarm leader is appointed from the central node;
- decentralized – the swarm leader is determined based on any algorithm and does not depend on the central control station;
- mixed – combines the advantages of centralized and decentralized strategies by selecting a swarm leader, if necessary, based on one of the algorithms with the transfer of control rights to the operator [10].

2. Determination of UAV parameters for group flight

For successful group autonomous flight of a group of UAVs, it is necessary to take into account: their density, the accuracy of positioning of each UAV, and the possibility and quality of communication with the control panel. An important condition is the stability of the UAV to wind gusts, turbulence, and other factors. Table 2 shows the characteristics of known quadcopter prototypes used for group flight.

A large number of quadcopters are required to create a UAV group, so it was decided to focus on the budget model CLOVER from COEX.

During the development of a quadcopter prototype based on the COEX CLOVER model, data available on the manufacturer's website was taken into account. Based on the prototype parameters and the xcopterCalc service [<https://www.ecalc.ch/xcopterCalc.php>], calculations were performed to obtain realistic results regarding flight time, compatibility, and the correctness of the selected components. Thanks to the xcopterCalc component database, it is possible to vary the parameters to obtain optimal values.

Table 2

Quadcopter prototypes used for group flights

Model name	Weight, g	Brightness	Battery life, min	Positioning technology	Network connection (Wi-Fi)	Appearance
IFO	786 (without battery)	Huge LED light	25	GPS RTK	2.4 GHz and 5.8 GHz range	
EMO	540	12 lamps	35	GPS RTK	5.8 GHz WiFi	
CLOVER	1000	LED strip	15	GPS	2.4 GHz	

The service has a wide range of settings and a large database of electronic components. The main settings that must be specified for the calculation (Fig. 3):

- model weight. When selecting a screw motor unit, the calculator automatically takes into account their weight. Considering the maximum take-off weight of 1000 g specified in the prototype, let's consider a model weighing 1200 g, including the screw motor unit;
- number of screws – for the model under development, as in the prototype, for clarity, the service immediately draws the resulting configuration;
- frame dimensions – the diagonal distance between the motor axes, the size is specified as 450 mm;
- battery (lead-acid battery) – «LiPo 4500 mAh 35/50C» battery and «Nominal» charge status; the number of batteries connected in parallel is specified in the «P» field (if using a circuit with several batteries);
- regulator (ESC) – select a speed regulator from the list based on your needs – max 60A;
- in the attachments, specify the total consumption and weight of all additional equipment to be installed, for example: camera, flashlight, etc.;
- motor – select based on the intended use of the multicopter, SunnySky V4006-380;
- propeller – for the initial calculation, we will take propellers from the popular manufacturer DJI with a diameter of 254 mm, three blades, and a pitch of 127 mm.

General Model Weight: 1200 g incl. Drive 42.3 oz # of Rotors: 4 flat Frame Size: 450 mm 17.72 inch FCU Tilt Limit: no limit Field Elevation: 113 m ASL 371 ft ASL Air Temperature: 25 °C 77 °F Pressure (QNH): 1013 hPa 29.91 inHg	
Battery Cell Type (Cont. / max. C) - charge state: LiPo 4500mAh - 35/50C normal Configuration: 4 S 1 P Cell Capacity: 4500 mAh 4500 mAh total max. discharge: 85% Resistance: 0.0036 Ohm Voltage: 3.7 V C-Rate: 35 C cont. 50 C max Weight: 116 g 4.1 oz	
Controller Type: max 60A Current: 60 A cont. 60 A max Resistance: 0.0045 Ohm Weight: 80 g 2.8 oz Accessories: 0 A Current drain: 0 A Weight: 0 g 0 oz	
Motor Manufacturer - Type (Kv) - Cooling: SunnySky good - V4006-380 (380) search... KV (w/o torque): 380 rpm/V no-load Current: 0.4 A @ 10 V Limit (up to 15s): 450 W Resistance: 0.17 Ohm Case Length: 18 mm 0.71 inch # mag. Poles: 24 Weight: 68 g 2.4 oz	
Propeller Type - yoke twist: DJI 0° Diameter: 10 inch 254 mm Pitch: 5 inch 127 mm # Blades: 3 PConst / TConst: 1.10 / 1.0 Gear Ratio: 1 : 1 calculate	

Fig. 3. Calculation parameters

Figure 4 shows the results and configuration of the multicopter. It is important to understand that the calculation is approximate and does not give guaranteed results, but it clearly shows how flight characteristics change when components and other parameters are changed. The figure shows:

- motor efficiency at optimal and maximum operating modes and in hover mode;
- propeller motor group efficiency in hover and maximum operation;
- payload, maximum roll, rate of climb, and probable time with flight range.

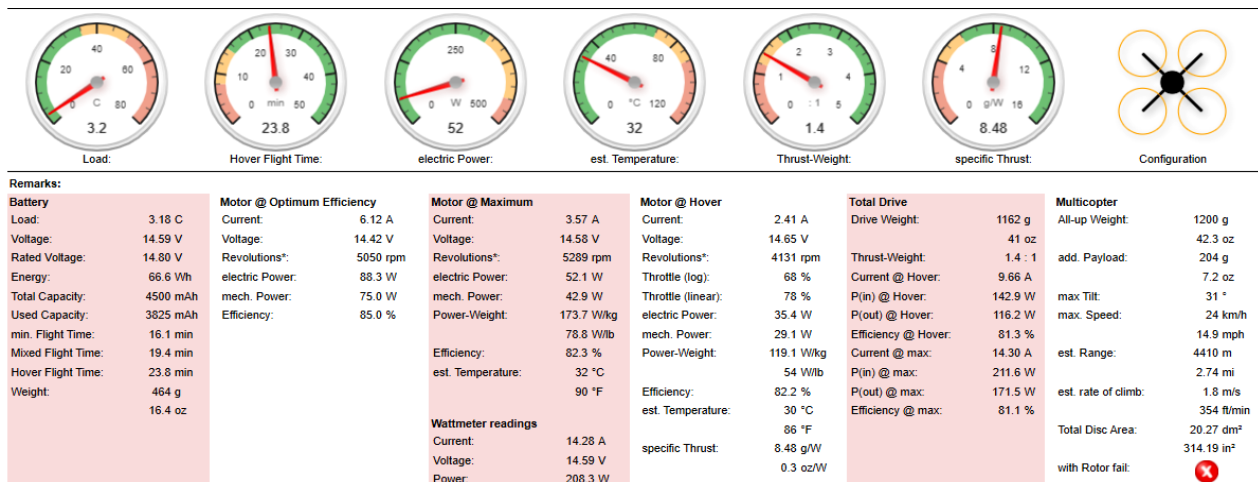


Fig 4. Calculation results

Let us consider the results of calculating the flight range using the graph shown in Figure 5. The left side of the graph shows the flight time (in minutes), the bottom side shows the speed of the UAV during flight (in kilometers per hour), and the right side shows the flight range (in meters). The green line shows the best range of flight time and distance data obtained. Depending on the weather, like if there's air resistance or not, there are two types of time and range.

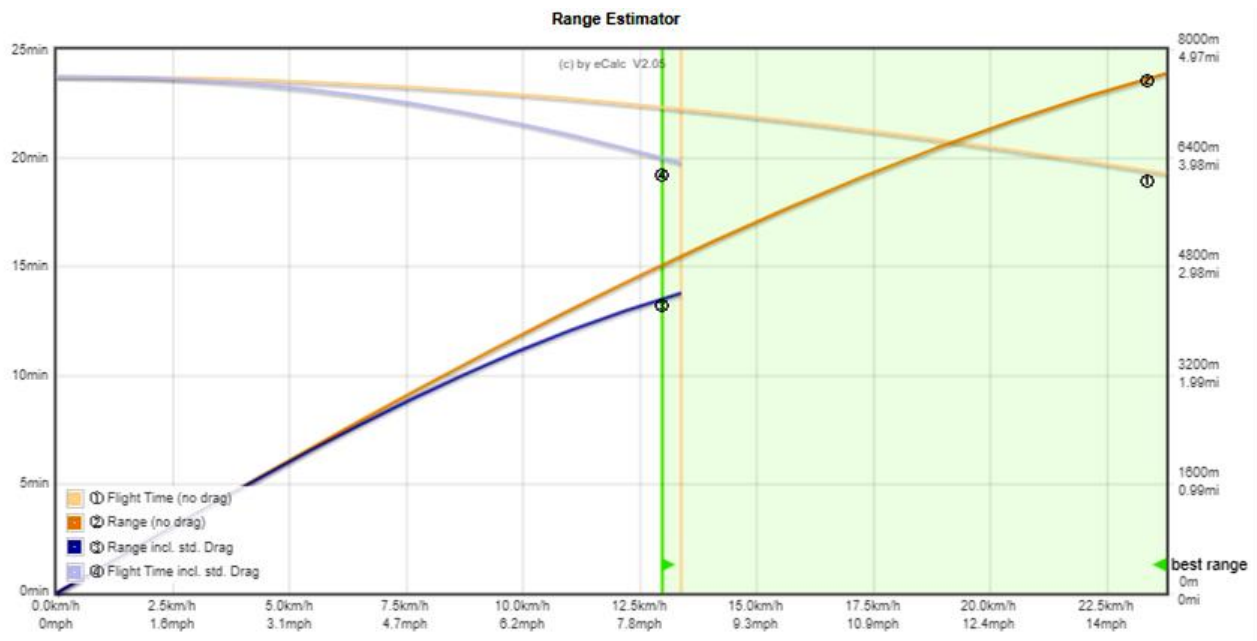


Fig. 5. Simulated graph of flight time and distance

Figure 6 shows the characteristics of the motor. The total power is marked with a circle on each characteristic curve. Based on the graph, we can see that the electric power limit of the motor can be even higher. The maximum flight time shown in Figure 5 (23.8 min) is calculated for the planned situation with the battery discharged to 85%. In reality, the flight time to avoid complete battery discharge will be no more than 19 minutes, i.e., a maximum discharge of 80%, which is sufficient for a group flight of UAVs.

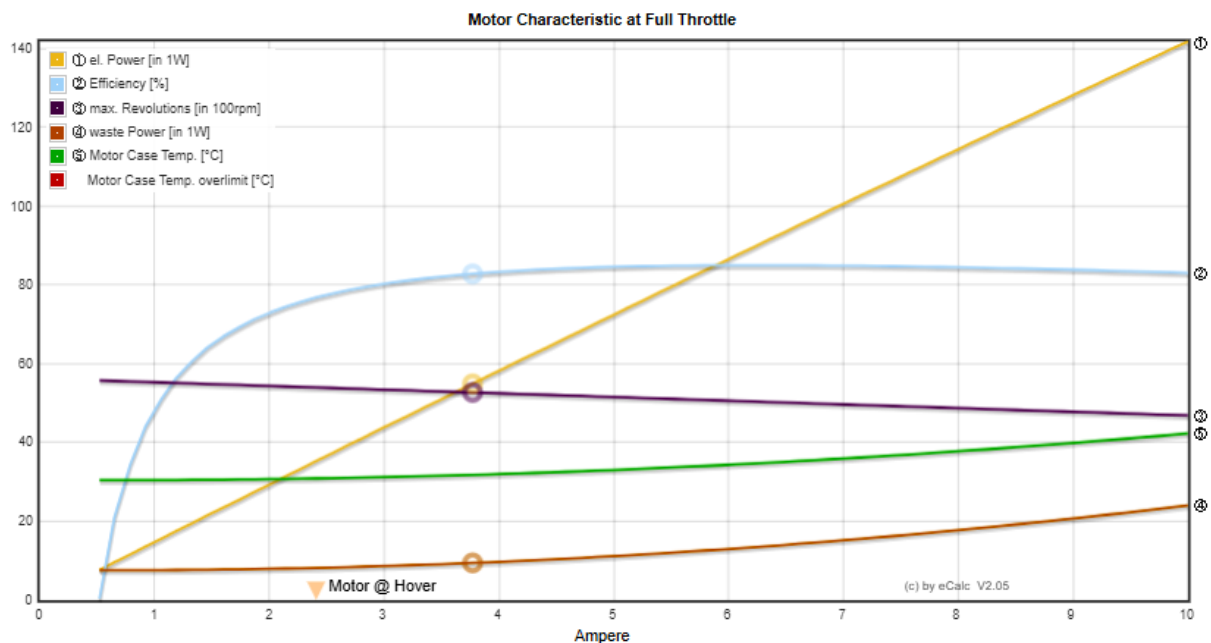


Fig. 6. Engine characteristics at full throttle

Considering that the calculation accuracy is $\pm 15\%$ and the UAV does not maintain a constant speed, the data obtained will differ slightly from the real scenario.

3. Mathematical models for constructing airspace maps and assessing conflicts

The increase in the number of small unmanned aerial vehicles has led to the emergence of a new task of unmanned traffic management (UTM). The main task of UTM is to estimate the throughput capacity – how many UAVs can be safely deployed with the possibility of successful management in a given airspace (AS). This task can be solved by estimating the following parameters:

- factors limiting AS capabilities, such as the emergence of intractable conflicts (if their probability is high, then measures to manage them should be identified);
- excessive UAV noise;
- interference with the UAV operator's communications (taking into account cybersecurity requirements, as encryption protocols require higher bandwidth).

The AS UTM throughput capacity assessment is based on air traffic management models and methods [11, 12], which mainly deal with flight planning from airport to airport. The UTM model differs in that it involves a larger number of aircraft and operators, a variety of flight tasks, and the ability to take off and land on unprepared sites. The limitations of the UTM model are the consideration of permitted movement zones for UAVs and restrictions on height above ground level (150 meters) [13].

Let us consider the stages of introducing UAVs into a single airspace. Based on NASA recommendations, four main stages have been identified, corresponding to levels of technological capability:

- 1) UAV motion control technologies for agricultural monitoring, firefighting, and infrastructure monitoring operations;
- 2) technologies for dynamic correction of airspace availability and management of unforeseen circumstances;
- 3) technologies for determining safe separation between private UAVs in moderately populated areas;
- 4) technologies for UAV operation in denser urban areas for tasks such as information gathering and package delivery.

Let us consider the objectives of these stages.

Stage 1. Testing of single UAV flights. Studying the demand for flights in sparsely populated areas.

Stage 2. Initial integration of UAVs into non-aggregated airspace, introduction of a basic UTM model.

Stage 3. Integration of unmanned aircraft systems (UAS) into the airspace of piloted aircraft, introduction of regional UTM.

Stage 4. Transparency in UAS management and situational awareness of all air traffic participants, creation of a centralized UTM.

Various information subsystems are used to determine flight-permitted areas and assess current traffic [14]. For example, the US Federal Aviation Administration website [15] provides maps showing flight-permitted areas, areas with flight altitude restrictions, and the possibility to register and obtain a flight license in electronic form.

To create such a subsystem in Ukraine, it is necessary to assess the probable traffic and collect information about no-fly zones.

The concept of the LiU UAV traffic probability map model is based on the following works:

- Dutch model «Metropolis» [16] – a probability model where aircraft are distributed evenly in a given AS. In general, the flight direction of UAVs is evenly distributed in a given AS circle (assuming that UAVs fly at the same altitude, i.e., the model is two-dimensional). This allows the probability of conflicts to be obtained;
- the Cal model [17] improves the previous model in terms of selecting flight end points. They are selected based on population density. In this case, each UAV flight is a line connecting the starting and ending points, raised to a given altitude.

Due to the rational selection of the throughput capacity of a given AS, limited by the number of UAVs, at which their safety decreases due to the emergence of conflicts involving three or more UAVs, the quality of UTM is improved.

To determine the intensity of UAV traffic, modeling methods are used with the application of a Poisson process model and a model for constructing and evaluating a random geometric graph.

The following input data are required to **build a map of** the studied AS for UAVs:

- region of interest R ;
- flight intensity $D(g)$, given for each point $g \in R$;
- observation duration T ($T = 12$ hours, i.e., daily traffic);
- the expected number of UAV operations during the time T (the parameter N changes during the experiments).

Considering the extremely low altitudes of AS, we assume that the demand for AS is generated according to the Cal model: the start time of a flight from

any point $a \in R$ is formed according to Poisson's law, according to which the intensity is proportional to the intensity of flights to a given point:

$$\lambda_s(a) = \frac{N}{T} \frac{D(a)}{\int_R D dA}. \quad (1)$$

Point b is randomly selected as the end point of the flight. Based on the given intensity of flights, the probability that the flight will end at the point $b \in R$ is

$$p(b) = \frac{D(b)}{\int_R D dA}. \quad (2)$$

The developed model calculates the point distribution of traffic. The intensity of the appearance of a UAV flying from a to b at the point $g \in ab$ is equal to the Poisson process:

$$\lambda_{ab} = \lambda_s(a) p(b) = \frac{N}{T} \frac{D(a)D(b)}{\left(\int_R D dA\right)^2}. \quad (3)$$

Integration for all origin-destination pairs allows us to determine the intensity of UAVs at any point g :

$$\lambda(g) = \int_{ab \ni g} \lambda_{ab} dA dB = \frac{1}{\left(\int_R D dA\right)^2} \frac{N}{T} \int_{ab \ni g} D(a)D(b) dA dB, \quad (4)$$

where $ab \ni g = \{(a, b) \in R^2 : g \in ab\}$ is the set of endpoints of all segments containing g .

To create a simulation model, the above model must be represented discretely: on R , we fix the grid L and the flight intensity $D(g)$, which is set for each grid cell $g \in L$. The start time of flights from any cell of the grid $a \in L$ forms a Poisson process, the intensity of which is

$$\lambda_s(a) = \frac{N}{T} \frac{D(a)}{\sum_{x \in L} D(x)} \quad (5)$$

is proportional to the intensity of flights in the cell, and the target flight cell b is randomly selected based on the density of the same probability:

$$p(b) = \frac{D(b)}{\sum_{x \in L} D(x)}. \quad (6)$$

The direct path of the UAV between cells a and b is represented by a sequence of grid cells through which the UAV passes (Fig. 7). Let's assume that the UAV spends the same time $t = l/v$ in each cell of the path, where l is the length of the cell side, v is the UAV speed.

Similarly to the continuous case for any cell $g \in ab$, UAVs flying from a to b enter cell g according to a Poisson process with intensity:

$$\lambda_{ab} = \lambda_s(a) p(b) = \frac{N}{T} \frac{D(a)D(b)}{\sum_{x \in L} D(x)}. \quad (7)$$

Summing up all the origin-destination pairs, we get that in general, the UAV is located in any cell g of the segment ab with intensity

$$\lambda(g) = \sum_{ab \ni g} \lambda_{ab} = \frac{1}{\left(\sum_{x \in L} D(x)\right)^2} \frac{N}{T} \sum_{ab \ni g} D(a)D(b), \quad (8)$$

where $ab \ni g = \{(a, b) \in L^2 : g \in ab\}$ is the set of endpoints of all segments containing g .

We call the graph of the function

$$m(g) = \sum_{ab \ni g} D(a)D(b) \quad (9)$$

the UAV map, which shows the probability of UAVs appearing in different cells. In general, the number of UAVs $n(g)$ that can be seen simultaneously in cell g at time t is calculated as follows:

$$\overline{\lambda(g)} = \lambda(g)t = \frac{t}{T \left(\sum_{x \in L} D(x)\right)^2} Nm(g). \quad (10)$$

In the experiments, the values of t and T do not change. Therefore, to simplify the formulas, we normalize the density to.

$$\frac{t}{T \left(\sum_{x \in L} D(x)\right)^2} = 1, \quad (11)$$

and, respectively

$$\overline{\lambda(g)} = Nm(g) \quad (12)$$

We will **evaluate UAV conflicts** using a random geometric graph model.

The Random geometric graphs (RGG) model will be used to model the number of UAV conflicts and collisions within a given traffic volume. A RGG is a graph $G(S, R)$ with a set of vertices S , which is obtained by placing n vertices randomly on the plane, with two vertices connected by an edge if the Euclidean distance between them is no more than R . Given a number $k > 0$, where $p_k(S, R)$ denotes the probability that $G(S, R)$ has a connectedness of size at least k :

- $p_1(S, R) = 1$, since any vertex is a connected component of size 1;
- $p_2(S, R)$ – the probability that $G(S, R)$ has an edge;

- $p_n(S, R)$ – the probability that the graph is connected;
- $p_k(S, R) = 0$ at $k > n$, since $G(S, R)$ has only n vertices.

Suppose that R is a conflict event, e.g., two UAVs at a distance R may collide. Then, for n randomly distributed UAVs (Fig. 7), the existence of a subnetwork with connection size k in $G(n, r)$ means that k UAVs are in conflict (Fig. 8).

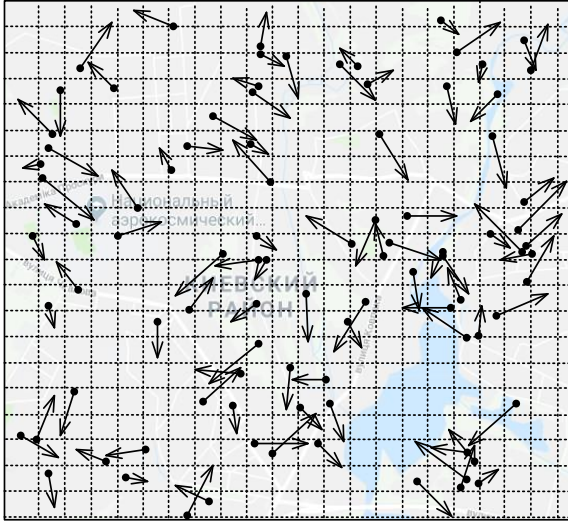


Fig. 7. UAV location map

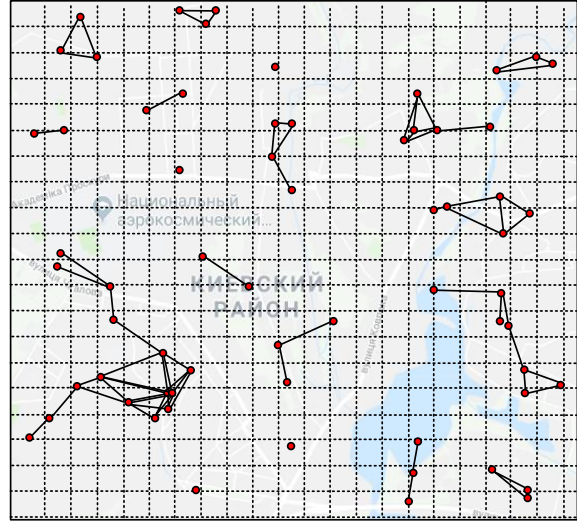


Fig. 8. Random geometric graph

With a small subnetwork size ($k = 2$), the conflict can be resolved using simple rules (e.g., the right maneuver). For $k > 2$, a conflict can mean a security breach event or its probability $p_k(S, R)$.

The parameter R represents the technical capabilities of the UAV – communication quality, navigation accuracy, autopilot performance, etc. Therefore, for the practical implementation of the RGG model in the UTM system, it is necessary to find a distance R that will ensure a low probability $p_k(S, R)$. At the same time, it should be indicated what technical characteristics of the UAV can ensure this distance.

Consider the problem of estimating $p_k(N, R)$ – the probability of observing a connected component of size at least k in the daily traffic of intensity N . Fig. 9 shows a graph of modeling the probability $p_3(N, R)$ for $N = 10, 200000$ $R = 5, 300$ m.

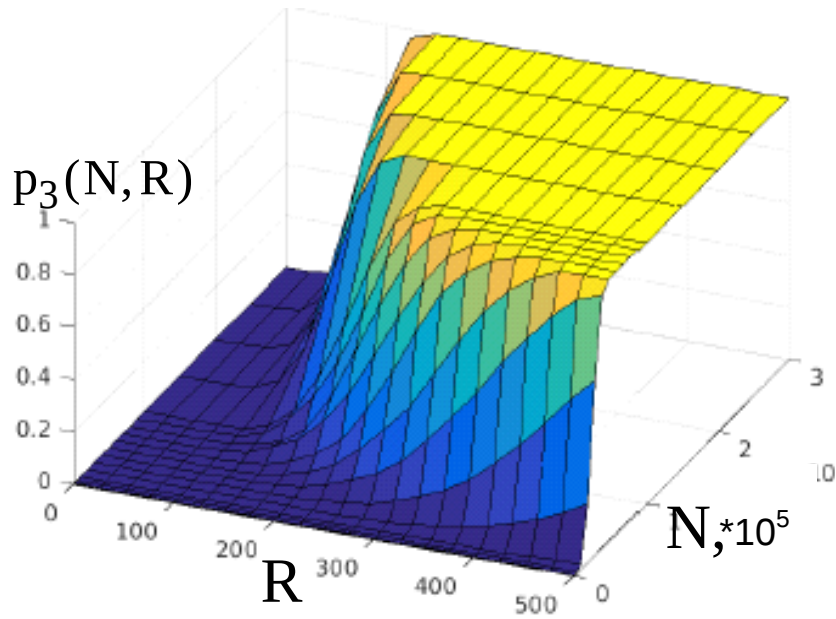


Fig. 9. Graph of probability modeling $p_3(N, R)$

Let's calculate the number of conflicts in RGG (related components) for two UAVs ($k=2$). The total expected number of such conflicts is the sum of conflicts for all pixel pairs:

$$C = \sum_{g, g' \in L^2} C(g, g') + \sum_{g \in L} C(g) \quad (13)$$

where $C(g, g')$ is the expected number of edges between nodes (UAVs) in cells g and g' ,

$C(g)$ – is the expected number of edges between UAVs in cell g .

Estimating $C(g, g')$ can be difficult because it depends on the exact location of the UAVs inside the cells g and g' (Fig. 10, a). For example, r and l may be in the following relationships with each other:

- if $r \gg l$ (Fig. 10, b), then this case can be ignored, since

$$C(g, g') = 0 \quad (14)$$

every time the cells g and g' are located «further than r » from each other and

$$C(g, g') = \frac{1}{2} \sum_{n, n'} n, n' \Pr\{n(g) = n\} \Pr\{n(g') = n'\}, \quad (15)$$

for the cells g and g' located «closer than r ». At $r = 500$ m and $l = 150$ m, this event does not occur, since the condition $r \gg l$ is not fulfilled;

- if $r \sim l$, you can reduce the grid so that the new cell size is $l' \ll r$, but calculations on the refined grid may take too much time;

– if $r \ll l$, edges between UAVs can exist only for neighboring cells g and g' . Moreover, the probability of such an edge is low (Fig. 10, c), since each UAV must fall into a band of r -width near the cell boundary with probability

$$(rl)/l^2 = r/l. \quad (16)$$

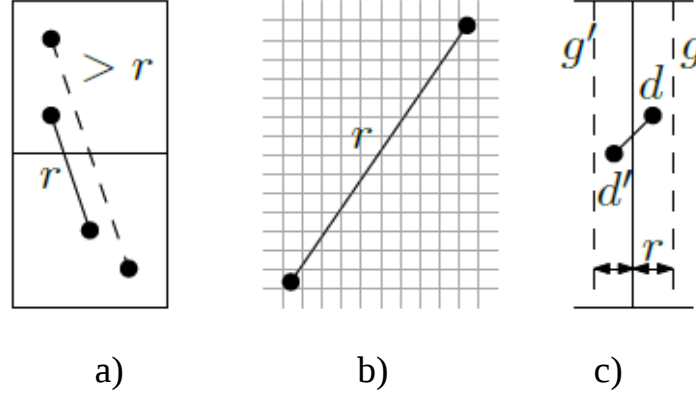


Fig. 10. Types of edges on the graph between nodes located in different cells

In addition, the distance between UAVs located near the cell boundary should be less than r (which gives an additional probability coefficient r/l).

In general, the probability of an edge between cells is $(r/l)^3$ (e.g., for $r = 5$ m, $l = 150$ m, it is about 10^{-4}) and, accordingly, the sum of edges between cells in (13) can be neglected.

Thus, we consider the sum of edges in the cells in formula (13) to obtain the expected number of conflicts. For a UAV at a given moment, the expected number of edges in cell g is $\frac{1}{2}n(n-1)\frac{\pi r^2}{l^2}$. Since the time of the UAV's movement in the cell is t , the total number of edges is

$$C(n) = \frac{1}{2}n(n-1)\frac{\pi r}{l} \quad (17)$$

Previously, it was believed that the number of UAVs in cell g ($n(g)$) according to the Poisson distribution law is $Nm(g)$. Thus, the expected number of edges in cell g is:

$$\begin{aligned} C(g) &= \sum_n C(n) \Pr\{n(g) = n\} = \frac{\pi r}{2l} \sum_n n(n-1) \Pr\{n(g) = n\} = \\ &= \frac{\pi r}{2l} \left(E[n^2(g)] - E[n(g)] \right) = \frac{\pi r N^2 m^2(g)}{2l}, \end{aligned} \quad (18)$$

where $E[n(g)] = \text{Var}[n(g)] = Nm(g)$.

When summing $C(g)$ over the entire grid during the entire simulation time, using the sum of discrete moments T/t , we obtain the expected number of conflicts during the flight:

$$C_d = \frac{\pi r N^2 T}{2lt} \sum_{g \in L} m^2(g). \quad (19)$$

Let's assume that a UAV collision occurs at $r = 5$ m. From the UAV map, we obtain the average flight duration:

$$\tau = \frac{1}{V|L|(|L|-1)\left(\sum_g D(g)\right)^2} \sum_{a,b \in L^2} |ab| D(a) D(b) \approx 2000c, \quad (20)$$

where $|L|$ is the total number of cells.

If we take into account the expected number of flight hours $H = N\tau$, we can see that this indicator will grow linearly, while the number of collisions C_d will grow quadratically. For example, for $N = 2457$, the expected number of conflicts is $C_d = 0.001$. Thus, it is possible to calculate the acceptable value of traffic intensity at which appropriate UTM measures are required.

4. Ontological model of flight zone categories

The probability of UAV conflicts in conditions of geometric variability of existing static obstacles, such as buildings and terrain, should be assessed taking into account the zones allowed for flight. To do this, we will create a model of no-fly zones.

To build a map of the zones allowed for flight, we built an ontological model of the regulatory document of the State Aviation Service of Ukraine «Temporary Procedure for the Use of the Airspace of Ukraine» [18].

The Protégé platform was chosen to build the ontology, which is a free, open ontology editor and a framework for building knowledge bases.

The structural elements of the ontology are arranged in a hierarchical order, with the relevant information needed at the lowest level of the hierarchy.

Let:

– a finite set of categories is given

$$C = (c_1, c_2, c_3, c_4),$$

where c_1 – forbidden zones,

c_2 – flight restriction zones,

c_3 – danger zones,

c_4 – temporarily reserved zones;

– a given feature space

$$P = P_1 \times P_2 \times \dots \times P_i \times \dots \times P_N,$$

where P_i is the set of features of the values of the i -th feature;

– is a finite set of elements of order

$$D = d_1, d_2, \dots, d_i, \dots, d_K.$$

– a feature function is given

$$f : D \rightarrow P, (fd_i) = (p_1, p_2, \dots, p_{Ni}),$$

is the feature description of the structural element d_i ;

– is an undefined function

$$\Phi : D \times C \rightarrow \{0,1\},$$

which for each pair (d_i, c_j) determines whether this element d_i , which has a feature description (fd_i) , belongs to the category c_j .

The following are the features used to define the categories:

$$[\text{airspace}] [\text{list of categories}] [\text{categorized accordingly}] [*] \rightarrow c_1$$

For d_i to be included in a category, a list of characteristics is needed to determine whether the category is a subclass or a concept:

The following are the features for defining concepts and subclasses of certain categories:

$$[*] [\text{flights are performed}] [\text{without} \mid \text{not}] [\text{list of attributes}] \rightarrow d_1$$

$$[*] [\text{flights are not performed}] [\text{list of signs}] \rightarrow d_1$$

$$[\text{Flight restriction areas}] [\text{list of features}] \rightarrow d_2$$

$$[\text{Dangerous areas}] [\text{list of signs}] \rightarrow d_3$$

where d_1 – signs of prohibited areas,

d_2 – signs of restricted areas,

d_3 – signs of hazardous areas.

Below are the signs for defining subclasses of categories:

$$[*] [\text{main streets}] [\text{list of features}] \rightarrow e_1$$

$$[*] [\text{operations zones}] [\text{list of features}] \rightarrow e_2$$

$$[*] [\text{roads of national importance}] [\text{list of features}] \rightarrow e_3$$

$$[*] [\text{designated objects}] [\text{list of features}] \rightarrow e_4$$

These concepts were divided into subclasses of categories because each concept describes one separate object, but they are connected by one common feature.

Based on the above features, the main categories and category subclasses were identified, on the basis of which the ontological model will be developed:

a) restricted areas:

- storage facilities for fuel, oil, gas, other hazardous substances and liquids, etc;
- penitentiary institutions;
- taxiways of airfields;
- places of accidents and disasters;
- power lines;
- pre-trial detention centers
- railways
- product pipelines;
- other important state and potentially dangerous facilities;
- power plants
- runways; and
- taxiways of heliports;
- state borders;
- industrial zones;
- sea ports;
- railway stations;
- central streets (cities, towns, villages);
- areas of operations (special, police, anti-terrorist);
- roads of national importance (international, national, regional, territorial);
- objects designated by state authorities (the Ministry of Defense of Ukraine,

the Ministry of Internal Affairs of Ukraine, the State Border Guard Service of Ukraine, the Security Service of Ukraine, the National Police of Ukraine, the National Guard of Ukraine, the State Protection Department, other military formations and law enforcement agencies established in accordance with the laws of Ukraine and in respect of which security is provided);

b) flight restriction zones:

- training grounds (bombing, firing, missile launches, use of military explosive devices, neutralization of ammunition by means of their explosion, landing)
- areas of firing to influence hydrometeorological processes in the atmosphere;

c) hazardous areas: the space above the high seas.

All categories and attributes are hierarchically linked by the IS-A relationship, as each higher-level attribute is a generalization of the lower level.

Based on these categories and concepts, an ontological model was built, which is shown in Fig. 11.

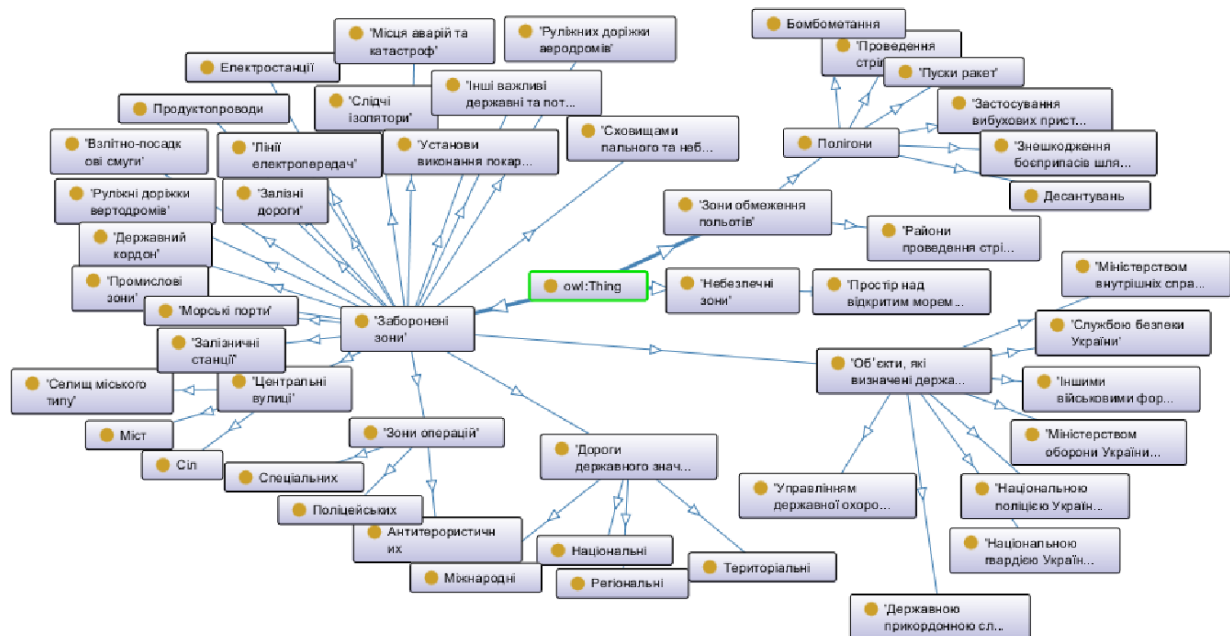


Fig. 11. Ontological model of flight zone categories

Using the developed ontological model, software was built that displays prohibited, restricted, and dangerous areas on the map. The basic scenario of the program's operation in a generalized form will consist of the following steps:

Step 1. Obtain data from the ontological model and save it to the program's knowledge base.

Step 2. Process the data using map queries. An important question is whether the existing electronic map can provide information on flight zones and how relevant this information will be. During the development of the software product, testing was only performed on available maps.

Step 3. Save the data.

Step 4. Output the data using an external service.

In this case, the external tools are Google API and a database server (Fig. 12).

The resulting map will enable the construction of a random geometric graph model for simulating the number of conflicts and collisions between UAVs within a given traffic volume, taking into account the areas permitted for flight. In addition, using the model developed in [10], it is possible to estimate the probability of UAV conflicts in a heavily built-up environment (Fig. 13).

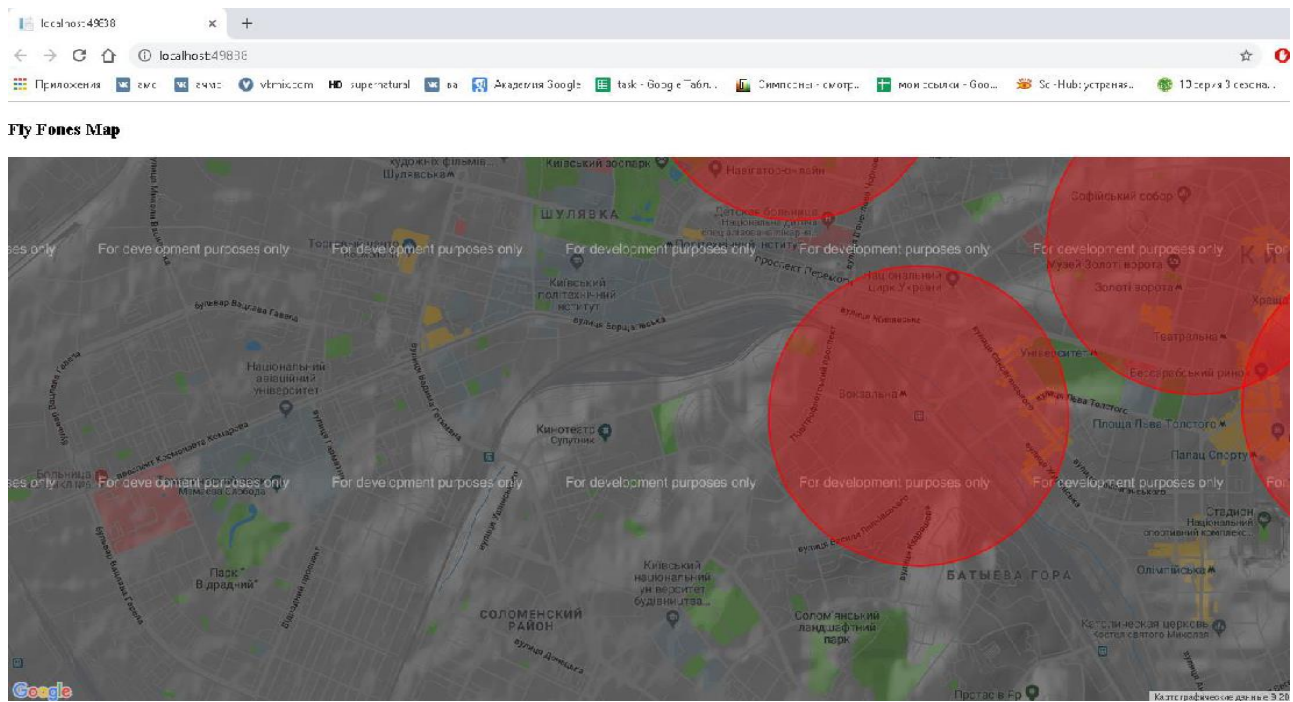


Fig. 12. Map showing flight zone categories

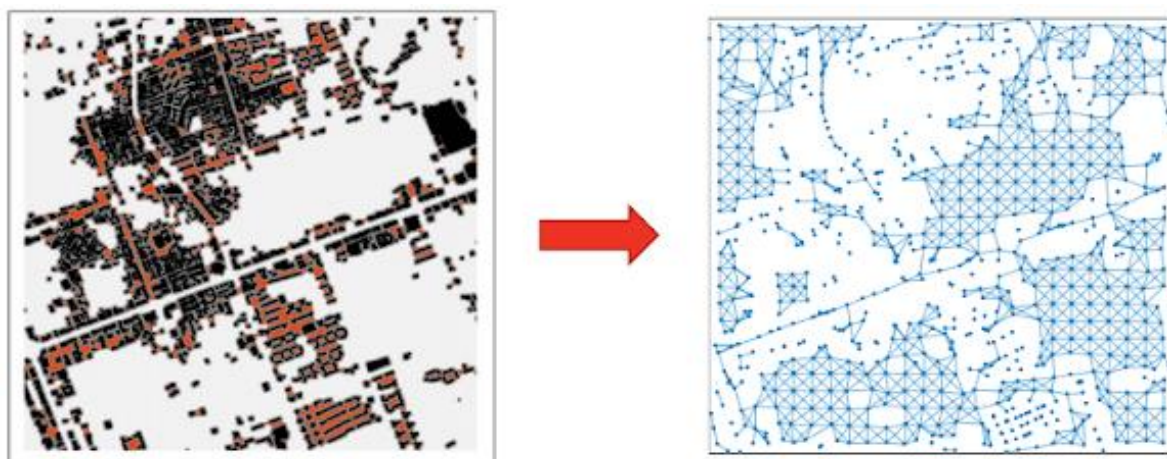


Fig. 13. Structuring of a heavily built-up urban environment

Conclusions

The levels of UAV autonomy and practical applications, including military applications at each level, are considered. In accordance with the classification of UAV autonomy levels, control levels are considered. UAV groups are classified based on the types of interactions between devices.

Based on data on the parameters of prototypes and the xcopterCalc service, calculations were performed to obtain realistic results regarding flight time, compatibility, and the correctness of the selected components. Thus, parameters were determined that ensure adaptive behavior in a swarm when environmental

conditions change. Graphical results of flight range and engine performance calculations are provided.

The parameters for assessing throughput capacity as the main task of unmanned traffic control are considered. According to NASA recommendations, four main stages corresponding to levels of technological capability have been identified. The concept of a probabilistic traffic map model for small UAVs has been applied. A generalized model of swarm interaction of UAVs has been constructed, which takes into account a multi-level control system and data exchange between devices. Modeling methods using a Poisson process model are used to determine the intensity of UAV traffic. A mathematical model of conflicts in airspace has been proposed, which allows predicting and avoiding collisions. UAV conflicts were assessed using a random geometric graph model.

An ontological model describes categories of airspace to support automated decision-making. Using the developed ontological model, software was created that displays prohibited, restricted, and dangerous areas on a map. The resulting map will enable the construction of a random geometric graph model for simulating the number of UAV conflicts and collisions within a given traffic volume, taking into account the zones permitted for flight.

The proposed parameterization method and mathematical and perception-oriented models contribute to improving the efficiency and safety of group flight of UAVs in a partially structured or changing air environment. The results obtained can be applied in the planning of military missions for the implementation of autonomous swarm strategies.

Thus, the work is aimed at determining the parameters of UAVs for group use, developing models for assessing air conflicts, and creating a conceptual basis for the implementation of a fully autonomous swarm. Establishing the relationships between the technical characteristics of UAVs, their behavior in a swarm, and the parameters of the air environment is the basis for further practical implementations.

їхньою поведінкою в рої та параметрами повітряного середовища є основою подальших практичних впроваджень.

The study was funded by the Ministry of Education and Science of Ukraine in the framework of the research project 0125U001544 on the topic «Methodology for ensuring the processes of monitoring and controlling the implementation of project and program portfolios for project offices in the context of Ukraine's reconstruction»

References

1. Orfanus, D., De Freitas, E. P., & Eliassen, F. (2016). Self-organization as a supporting paradigm for military UAV relay networks. *IEEE Communications Letters*, 20(4), 804–807. DOI: <https://doi.org/10.1109/LCOMM.2016.2524405>
2. Chen, W., Liu, J., & Guo, H. (2020). Achieving robust and efficient consensus for large-scale drone swarm. *IEEE Transactions on Vehicular Technology*, 69(12), 15867–15879. DOI: <https://doi.org/10.1109/TVT.2020.3036833>
3. Chen, W., Zhu, J., Liu, J., & Guo, H. (2024). A fast coordination approach for large-scale drone swarm. *Journal of Network and Computer Applications*, 221, 103769. DOI: <https://doi.org/10.1016/j.jnca.2023.103769>
4. Rumba, R., & Nikitenko, A. (2020). The wild west of drones: A review on autonomous-UAV traffic-management. *2020 International Conference on Unmanned Aircraft Systems (ICUAS)*, 1317–1322. DOI: <https://doi.org/10.1109/ICUAS48674.2020.9214031>
5. Campion, M., Ranganathan, P., & Faruque, S. (2018). UAV swarm communication and control architectures: A review. *Journal of Unmanned Vehicle Systems*, 7(2), 93–106. DOI: <https://doi.org/10.1139/juvs-2018-0009>
1. Shrestha, R., Oh, I., & Kim, S. (2021). A survey on operation concept, advancements, and challenging issues of urban air traffic management. *Frontiers in Future Transportation*, 2, 626935. DOI: <https://doi.org/10.3389/ffutr.2021.626935>
2. Hashim, H. A. (2025). Advances in UAV avionics systems architecture, classification and integration: A comprehensive review and future perspectives. *Results in Engineering*, 25, 103786. DOI: <https://doi.org/10.1016/j.rineng.2024.103786>
3. Druzhynin, E. A., Kovalevskyi, M. I., Pohudina, O. K., & Cheranovskiy, V. O. (2021). Methods and information technologies of implementation of unmanned aerial vehicles into the airspace of Ukraine. *Systems of Arms and military equipment*, 4(68), 84–90. DOI: <https://doi.org/10.30748/soivt.2021.68.12> (In Ukrainian)
4. Ryan, A., Zennaro, M., Howell, A., Sengupta, R., & Hedrick, J. K. (2004). An overview of emerging results in cooperative UAV control. *43rd IEEE Conference on Decision and Control (CDC)*, 1, 602–607. DOI: <https://doi.org/10.1109/CDC.2004.1428700>
5. Pohudina, O., Kritskiy, D., Koba, S., & Pohudin, A. (2020). Assessing unmanned traffic bandwidth. In Nechyporuk, M., Pavlikov, V., & Kritskiy, D. (Eds.), *Integrated Computer Technologies in Mechanical Engineering* (Vol. 1113). Springer. DOI: https://doi.org/10.1007/978-3-030-37618-5_38
6. Kopardekar, P., Rios, J., Prevot, T., Johnson, M., Jung, J., & Robinson, J. E. (2016, June). Unmanned aircraft system traffic management (UTM) concept of operations. *AIAA AVIATION Forum and Exposition*. DOI: <https://core.ac.uk/download/pdf/186445202.pdf>
7. Straubinger, A., Rothfeld, R., Shamiyeh, M., Büchter, K. D., Kaiser, J., & Plötner, K. O. (2020). An overview of current research and developments in urban air mobility – Setting the scene for UAM introduction. *Journal of Air Transport Management*, 87, 101852. DOI: <https://doi.org/10.1016/j.jairtraman.2020.101852>
8. Alharbi, A. A., Petrunin, I., & Panagiotakopoulos, D. (2022). Modeling and characterization of traffic flow patterns and identification of airspace density for UTM application. *IEEE Access*, 10, 130110–130134. DOI: <https://doi.org/10.1109/ACCESS.2022.3228828>

9. Ahmed, F., & Hawas, Y. E. (2015). An integrated real-time traffic signal system for transit signal priority, incident detection and congestion management. *Transportation Research Part C: Emerging Technologies*, 60, 52–76. DOI: <https://doi.org/10.1016/j.trc.2015.08.004>
10. Federal Aviation Administration. (n.d.). *United States Department of Transportation*. URL: <https://www.faa.gov>
11. Bereziat, D., Cafieri, S., & Vidosavljevic, A. (2022, December). Metropolis II: Centralised and strategical separation management of UAS in urban environment. In *12th SESAR Innovation Days*. URL: <https://enac.hal.science/hal-03945545/>
12. Sedov, L., & Polishchuk, V. (2018, June). Centralized and distributed UTM in layered airspace. In *8th International Conference on Research in Air Transportation* (pp. 1–8). URL: <https://undefined.github.io/publications/icrat.pdf>
13. Temporary procedure for the use of Ukrainian airspace (n.d.). URL: <https://ips.ligazakon.net/document/FN042940> (In Ukrainian)

CLOUD DEPLOYMENT STRATEGIES AND ARCHITECTURAL INNOVATIONS IN LEARNING MANAGEMENT SYSTEMS

Moiseienko A., Kuznetsova Yu.

The article provides a comprehensive analysis of modern technologies for deploying learning management systems (LMS) in cloud environments, with special attention to containerization, microservice architecture, and multi-tenant SaaS models. Practical examples of implementing Moodle, Canvas, Blackboard Learn, and Open edX in leading cloud platforms (AWS, Azure, GCP) are considered, including technical aspects of the infrastructure: scaling, ensuring fault tolerance, using Kubernetes, Docker, Terraform, CI/CD pipelines, and managed services. A comparative analysis of on-premise, SaaS, and containerized deployment models is conducted based on key criteria (flexibility, reliability, maintenance cost, scalability). Practical recommendations are offered for choosing the optimal LMS architecture depending on the implementation context and resources of the educational institution. The results of the study demonstrate the advantages of cloud technologies for developing educational web systems and increasing their availability, stability, and efficiency in the context of the digitalization of education.

Introduction

Over the past decade, the digital transformation of education has developed rapidly, which was made possible by the spread of modern information and communication technologies. One of the key innovations in this area was the introduction of learning management systems (LMS), which provide centralized organization, implementation, and monitoring of the educational process in an online environment [1, 2]. However, traditional LMSs, which operate on local servers of educational institutions, have a number of limitations related to scalability, reliability, maintenance costs, and the implementation of new functionalities [3, 4].

The search for solutions to overcome these problems has led to the active introduction of cloud technologies in the field of education. Cloud computing, in particular, the SaaS, PaaS, and IaaS models, open up new opportunities for flexible, secure, and efficient deployment of SNS in the environment of public or hybrid clouds [13]. This topic has become particularly relevant in the context of the COVID-19 pandemic and martial law in Ukraine, when the need for continuous access to educational resources has become critical, regardless of the geographical location or physical condition of the infrastructure of the educational institution [2, 4].

However, the practical implementation of a cloud-based CMS is accompanied by a number of challenges, from issues of data security and confidentiality to the difficulties of scaling and integration with other services. A particularly important problem is the optimal distribution of CMS components in a cloud environment, taking into account limited resources, performance requirements, fault tolerance, and cost [11, 12].

The object of research is modern learning management systems in the context of cloud technologies.

Subject of study – methods, models, and tools for implementing and optimizing cloud deployment of learning management systems, in particular, using microservice architecture, containerization, and formal optimization models.

The purpose of the study is the development, analysis, and generalization of effective approaches to implementing cloud-based learning systems with an emphasis on automated planning and resource optimization, which allows for flexible scaling and reliable operation of educational web systems.

To achieve the goal, it is necessary to solve the following tasks:

- systematize current approaches to cloud-based deployment of SNS;
- identify shortcomings and limitations of existing solutions;
- explore the potential of microservice and container architectures;
- to formulate a mathematical model of the problem of optimal placement of SNS in the cloud;
- to propose directions for further improvement of such systems.

Practical significance. The research results can be used by educational institutions, IT companies providing EdTech solutions, and researchers to improve the efficiency of cloud-based learning systems implementation. The proposed models contribute to reducing costs, increasing resilience to failures, and improving the user experience in online learning.

1. Overview of cloud deployment practices learning management systems

To analyze modern approaches to implementing and deploying learning management systems in cloud environments, it is necessary to consider practical examples of deploying popular LMS platforms in the infrastructures of leading cloud providers (AWS, Azure, GCP), as well as the architecture of multi-tenant SaaS solutions (multi-tenant) and containerized microservice deployments [5–7, 13].

As a result, the technical aspects of the architecture, deployment models, and the use of tools such as Kubernetes, Docker, Terraform, Helm, and DevOps pipelines for automation will be described in detail. Examples of the implementation of SNS with scalability, reliability, and performance in the cloud will be given, a comparative analysis of the models (on-premise vs. SaaS vs. containerized) will be performed in the form of a table, and conclusions will be formulated regarding the practical advantages and disadvantages of each approach [14].

1.1 Moodle in cloud environments AWS, Azure, GCP

Moodle on AWS is deployed using EC2 (web servers), Amazon RDS/Aurora (DB), ElastiCache (Redis), EFS (files), and ALB + CloudFront (load balancing and CDN). Auto Scaling groups provide dynamic scaling, and managed services increase reliability. Integration with AWS CodePipeline and CodeDeploy allows for CI/CD for continuous updates of Moodle [5].

Moodle deployment on Azure is based on ARM templates. The basic configuration uses a single VM, MySQL, and NFS. For larger workloads, scenarios with multiple web hosts, Redis cache, and GlusterFS for storage are used. Azure Database for MySQL simplifies administration, but compared to AWS, more components require manual management (VMs, file systems) [6].

Moodle on Google Cloud is implemented as a cloud-native solution: containerized in Docker, orchestrated in GKE (Kubernetes). Cloud SQL (DB), Filestore (files), Cloud Memorystore (Redis), Cloud Load Balancer, CDN and WAF (Cloud Armor) are used. The solution is fully managed and automated: no manual work with VMs, with CI/CD support via Cloud Build and Artifact Registry [7].

It can be concluded that all three platforms support scalable and reliable Moodle deployment. AWS provides the most out-of-the-box managed components. Azure is suitable for flexible scenarios, but requires more manual configuration. GCP offers the most modern cloud-native approach with deep containerization.

1.2 Architecture of the Canvas LMS multi-user SaaS solution

Canvas LMS [8] is a cloud-native platform with a multi-tenant SaaS architecture, where each client (university or school) is a separate «tenant» in a shared environment. Isolation is achieved by sharding the database: each client has its own shard, although all users work through a common web application.

The Canvas infrastructure runs on AWS and is built on Ruby on Rails for scalability. Inbound traffic passes through AWS WAF and is distributed to web

servers by Elastic Load Balancer, data is cached via Redis/Memcached. Data is stored in a DBMS cluster (Aurora or PostgreSQL), and exchange rate information is stored in S3 storage with delivery via CDN (CloudFront).

For background tasks (reports, conversions) queue servers are used that scale automatically. The platform is deployed in multiple AWS regions to meet data storage requirements. Administration, updates, scaling, and backups are handled by the provider (Instructure).

Canvas demonstrates a combination of multi-tenancy, horizontal scaling, and automated resource management that allows it to serve millions of users with a high level of reliability [8].

1.3 Managed services and integration with DevOps practices in Blackboard Learn SaaS

Serving over 100 million users, Blackboard Learn has evolved from an on-premise solution to a modern SaaS platform deployed in the AWS cloud. While historically the system was monolithic (Java), the SaaS version implements a multi-tenant architecture that allows educational institutions to operate in a single environment while maintaining data isolation [9].

The cloud infrastructure includes web hosts in multiple AWS availability zones, load balancers, caching (Redis/Memcached), and file storage based on Amazon S3 (in partnership with NetApp Cloud Volumes ONTAP). This allowed Blackboard to dynamically scale, including during the pandemic when the load increased by 3000%.

The databases have been migrated from Oracle/MSSQL to PostgreSQL on Amazon RDS, which has reduced costs, increased automation, and enabled the use of Read Replicas for analytics without compromising performance. The company is also reportedly using ElastiCache, CloudWatch, and AWS Lambda.

DevOps practices enable continuous integration and deployment (CI/CD). Blackboard implemented a fully automated release pipeline with containerization, infrastructure as code, and blue-green releases. Automated UI testing (powered by AWS Lambda) reduced testing time from days to seconds. With IaC, the company can scale resources in hours instead of weeks [9].

As a result, we can conclude that the platform is capable of serving millions of users with high SLAs, without downtime, with quick response to changes and updates «on the fly».

1.4 Cloud deployment of Open edX based on containerization and orchestration in Kubernetes

Open edX is a scalable online learning platform that includes a number of interconnected services: LMS, CMS (Studio), MySQL, MongoDB, ElasticSearch, analytics, etc. Previously, its deployment required complex configuration via Ansible, but today the community has moved to containerization and orchestration in Kubernetes [10].

Containerization of components. Each Open edX component is packaged in a separate Docker container using the «one service, one container» principle. This simplifies deployment, provides isolation, scalability, and reduces dependencies between services. Persistent volumes are used for stateful components [10].

Orchestration in Kubernetes. In production environments, Open edX is deployed as a set of pods in a Kubernetes cluster. YAML specifications are created for the LMS, CMS, DB, and other services, specifying environment variables, recovery policies, and state retention (via PersistentVolumeClaim). For simplicity, an official Helm Chart is used to automate the deployment of the entire platform [10].

Scalability and continuous operation. Kubernetes provides horizontal and vertical auto-scaling. For example, during a large course, the number of LMS pods automatically increases. Rolling updates are also supported, which allows you to update the system without downtime. The EFK stack and probes (liveness/readiness) are used for monitoring and logging [10].

Safety and isolation. Containers provide service isolation. Kubernetes allows you to restrict network communication through Network Policies. This is critical for multi-tenant deployments, where different clients use separate instances on a shared cluster [10].

Therefore, the modernized approach to deploying Open edX significantly increases the flexibility, scalability, and reliability of the system. The use of Kubernetes and Helm ensures efficient management of a multi-component LMS and makes it suitable for high loads and dynamic usage scenarios.

2. Comparative analysis of cloud implementation models learning management systems

To summarize, let's look at the key characteristics of different approaches to deploying learning management systems – from traditional local (on-premise) to modern multi-cloud. Table 1 compares three basic models [11, 14]:

1. Local/self-contained deployment – when the institution itself installs the LMS on its servers or rented VMs (monotenant architecture);

2. Cloud SaaS (multi-tenant) – when the LMS is provided as an online service by the provider and serves many clients in one shared environment;

3. Cloud containerized deployment – when the institution or provider deploys the LMS in a cloud infrastructure using containerization and microservices (can be either a separate instance for one institution or multi-instance on a shared cluster).

This classification covers most of the practices described above, i.e. Moodle traditionally implements the first and third models (standalone deployment on VM or modern containerized on GCP) [5–7], Canvas and Blackboard represent the second model (SaaS) [8, 9], and Open edX is used in all three variants (there are institutions that install Open edX locally; there are SaaS providers such as EdX/Open edX; there are containers in the cloud) [10]. The main comparison criteria: scalability, support requirements, financial model, customization capabilities, and reliability/fault tolerance are given in Table 1.

Also, it should be noted that each of the models has its own optimal use scenarios. On-premises deployment is appropriate where complete autonomy and control over data are critical (for example, in military or closed corporate networks). The SaaS model is best suited for a quick start and minimizing technical risks, which is especially important for small institutions with limited IT staff. Containerized cloud deployment, in turn, allows for a combination of scalability and control, which makes it advantageous for universities with large audiences or for national educational platforms.

Additionally, security and regulatory compliance (GDPR, FERPA, etc.) are important. In a SaaS model, much of the responsibility lies with the provider, while in on-premise and containerized systems, the institution independently defines access and security policies. This can be a decisive factor in choosing an architecture.

Equally important is the issue of integration with other services: on-premise solutions are often more difficult to integrate with modern analytical tools or video conferencing systems, while SaaS and Kubernetes-based models usually offer ready-made APIs and modules for interaction. Finally, one should consider the risk of «vendor lock-in» – dependence on a single cloud service provider, which can limit an institution's flexibility in the future.

Table 1

Comparison of SSN deployment models

Criterion	Local / self-deployment (on-premise or own VMs)	Cloud SaaS (multi-tenant model)	Containerized cloud deployment (cloud-native microservices)
1	2	3	4
Scalability	Limited by local resources; scaling requires purchasing and manually configuring new hardware or VMs. High peak loads are difficult to handle without redundancy.	Almost unlimited horizontal scalability due to provider resources; servers/capacity are added dynamically according to demand (auto-scaling in the cloud). Customers do not notice the increase in load - this is the provider's problem.	Flexible scalability thanks to Kubernetes and cloud services: the system can automatically scale individual components (pods) based on load metrics. It is easy to scale different services independently (for example, only LMS web servers or only DB by increasing resources or shards).
Costs and payment model	Large capital costs at the beginning: purchase of servers or long-term lease of VMs, purchase of software licenses. Ongoing costs for electricity, cooling, IT staff. Scaling requires additional investments.	OPEX model: predictable subscription or payment per user/year. Low initial costs (no need to buy equipment). The provider distributes infrastructure costs among many customers, which reduces TCO for each. Scaling means subscription growth, but does not require direct customer action.	Pay-as-you-go. No capital costs for servers, but DevOps team costs. Optimal resource utilization through auto-scaling can make maintenance cheaper (minimum resources during off-hours). However, the complexity of the platform may require experts, which is also worth it.
Support and updates	The responsibility lies entirely with the institution's IT department: they must independently apply patches, update the LMS to new versions, and monitor server security. Downtime is possible during updates.	All updates are performed by the SaaS provider, often continuously (Continuous Delivery) and with minimal disruption. The customer always receives the latest version without any effort. There is no need to have a large IT staff to support it – «LMS as a Service».	Automation greatly simplifies support: CI/CD pipelines can deploy new versions of containers themselves. Kubernetes allows rolling updates without downtime. But the responsibility for configuring updates lies with the team that implemented the system. Requires high competence of DevOps engineers to maintain the infrastructure.

Continuation of the table 1

1	2	3	4
Customization and control	Maximum customization: the institution has full control over the servers and software, can modify the LMS source code, install non-standard plugins, configurations, integrate with internal systems at a deep level. This is important for organizations with special needs (e.g., specific authentication).	Limited flexibility: A multi-user SaaS system typically does not allow for customizing the source code for a specific client. Customization is limited to configuration options and adding supported plugins. Deep integrations are only possible through the provided APIs. In return, the client gets stability and proven solutions, no «manual» intervention.	High flexibility: the system is deployed under the control of the customer's (or provider's) team, you can modify the container configuration, add your own services (e.g. additional analytical modules), choose component versions. A Kubernetes-oriented system easily integrates with other cloud services (e.g., you can add an external analytics system, Lambda server, etc.). Customization is almost as broad as on-premise, but requires skills.
Reliability and scale	Depends on redundancy in the local infrastructure. High reliability can be achieved, but it is expensive: server duplication, database clustering, backup power. For many institutions, a typical configuration is one application server and one database server, which are points of failure. Fault tolerance and disaster recovery are entirely the responsibility of the institution.	Deployed on a professional infrastructure with redundancy at all levels (multiple datacenters, database replication, daily backups, etc.). The provider guarantees an SLA (usually 99.9% and above uptime). The scale is global: users around the world receive acceptable performance through CDN and geo-distribution. On the other hand, a general SaaS failure (although unlikely) affects all customers at once.	Built on cloud capabilities, such a system can be very reliable with the right configuration: multi-AZ deployment, use of managed databases with automatic failover, regular backups, etc. Kubernetes itself provides service recovery in the event of container or node failure. Disaster recovery and reserves on the customer side: you need to think about, for example, a backup cluster or a backup strategy in the cloud. Scaling to very large loads is possible, but you need to follow the limits of the cloud provider and optimize the architecture (for example, sharding the database when growing, distributing components across microservices).

3. Analysis of cloud deployment models learning management systems

The review showed that there is no clear universal model for deploying a cloud computing system – each approach has its strengths and weaknesses (summarized in Table 1), and the choice depends on the needs and capabilities of a particular organization. Here are some practical conclusions regarding the main models of cloud deployment.

1. Traditional local deployment (on-premise/self-hosted)

The main advantages include: full control over the system, data, and update process; the ability to deeply customize the LMS to your requirements [10]; and no dependence on external providers (relevant for closed networks or strict security policies).

Disadvantages: high costs for infrastructure and IT personnel – investments in servers, their maintenance, and backup are necessary; limited scalability – it is difficult to respond quickly to a sharp increase in the number of users or load (you need to purchase enough capacity in advance); risk of downtime and difficulty in ensuring high availability – the administrator is responsible for backup, fault tolerance, updates without stopping the system, etc. In modern conditions, an on-premise solution is justified only for small institutions with very specific requirements or where the use of the cloud is impossible for security reasons [11].

2. Multi-user cloud SaaS solution

The main advantages are minimal hassle for the customer – the deployed system is available immediately after signing the contract, all aspects of technical support (updates, patches, monitoring) are provided by the provider; high reliability and scalability «out of the box» – the provider provides the necessary resources, geographical distribution, and backup, which would be very expensive at individual cost; budget predictability – subscription payment that correlates with the number of users or other metrics, without capital expenditures.

Disadvantages: less flexibility and customization options – you mostly have to work within the framework of standard functionality, since the common platform does not allow you to change it for only one client (it is impossible to make changes to the LMS code if they are not provided for everyone); dependence on the provider – both technical (in the event of a service failure, the client will not fix anything himself, he must wait for a support response) and strategic (changes in prices, policies, or termination of product support by the provider directly affect the institution). Also, the issue of data security requires trust in the provider: educational

data is stored in the cloud, and the institution must be confident in its security and compliance with regulatory requirements (GDPR, FERPA, etc.).

In general, the SaaS model is optimal for most cases, especially when the priority is rapid implementation and reduced IT costs [5-9].

3. *Containerized cloud deployment (microservice architecture in IaaS/PaaS)* combines the advantages of the cloud and independent control

The customer (or the cloud operator providing the service) can flexibly configure the system using modern DevOps tools: automate deployments, scale components in real time, and optimally use resources.

Containerization allows you to isolate LMS services, which simplifies updates and ensures stability (the failure of one component will not «bring down» the entire system). Kubernetes orchestration brings self-healing to the system (automatic restart of crashed containers) and makes it easier to achieve high availability without manual intervention. In addition, the cloud-native approach provides portability: you can relatively easily move containers to another cloud platform or deploy a hybrid solution (part of the services run locally, part in the cloud) – this reduces the risk of being tied to a single vendor.

Disadvantages: significant complexity of implementation – experts in containerization, CI/CD setup, cluster security, etc., are required; the entry threshold for small teams is high. In fact, deploying a Kubernetes cluster with Open edX or Moodle is within the power of only a fairly mature team, while small organizations are easier to choose SaaS or use third-party integrators. It should also be noted that migration to microservices does not completely eliminate the problems of database scaling or transactional integrity – with a very high load, you will still have to implement database sharding or distribute services across multiple clusters (which requires additional efforts in architecture design). Therefore, the container model is most effective for medium and large implementations, where there are resources to support it and the need for flexibility/scale, and is less justified for small installations [11–14].

Conclusions

Cloud technologies have opened up new horizons for learning management systems, enabling rapid scaling to any audience and high availability of educational platforms around the world.

The practices discussed in this section – from hosting Moodle on AWS/Azure/GCP to SaaS architectures of Canvas and Blackboard and containerization

of Open edX – demonstrate that the competent use of cloud services and DevOps approaches can significantly improve the effectiveness of an LMS.

In particular, the transition to a SaaS model allows you to focus on the educational process, leaving the technical aspects to the provider, while the microservices approach provides maximum flexibility and the potential for optimization to your needs. The choice of model depends on the context: for small institutions, SaaS is often optimal (minimum hassle and high reliability), for large universities or consortia with their own IT resources - containerized deployment can provide the desired control and economies of scale, and some specific scenarios may still require on-premise (for example, military academies or remote closed networks).

In general, the global trend is clearly moving towards the cloud: cloud implementation and deployment technologies for E-learning are becoming an integral part of the development of educational web systems, ensuring their readiness for the challenges of mass and continuous learning in the digital age.

Prospects for further work in this direction are associated with the development of adaptive and intelligent LMSs that will use artificial intelligence methods to personalize the learning process. An important vector will be the use of hybrid architectures that combine local resources and public clouds to ensure greater security and reliability. In addition, the implementation of unified interoperability standards is promising, which will simplify the integration of LMSs with external educational services.

References

1. Al-Busaidi KA, Al-Shihi H. Key factors to learning management system implementation success: A case study of Sultan Qaboos University // *Education and Information Technologies*. – 2019. – Vol. 24, No. 2. – P. 1313–1335. – DOI: 10.1007/s10639-018-9831-y
2. Pappas IO, Pateli AG, Giannakos MN, Chrissikopoulos V. The interplay of online learning motivations and cloud-based LMS quality on students' satisfaction // *Computers & Education*. – 2019. – Vol. 132. – P. 75–92. – DOI: 10.1016/j.compedu.2019.01.004
3. Kumar A., Sharma R. Cloud computing for education: A modern approach to teaching and learning // *Journal of Cloud Computing: Advances, Systems and Applications*. – 2021. – Vol. 10, No. 15. – P. 1–17. – DOI: 10.1186/s13677-021-00234-6
4. Singh G., Hardaker G. Cloud-based learning management systems: A review on features and architecture // *Education and Information Technologies*. – 2020. - Vol. 25. – P. 5131–5154. – DOI: 10.1007/s10639-020-10234-7
5. Amazon Web Services. Deploying Moodle on AWS [Electronic resource]. – Access mode: <https://aws.amazon.com/quickstart/architecture/moodle/> (date of application: 08/23/2025).
6. Microsoft Azure. Deploy Moodle on Azure [Electronic resource]. – Access mode: <https://azuremarketplace.microsoft.com/en-us/marketplace/apps/bitnami.moodle> (date of application: 08/23/2025).

7. Google Cloud Platform. Moodle on GCP using GKE [Electronic resource]. – Access mode: <https://cloud.google.com/solutions/moodle> (date of application: 08/23/2025).
8. Instructure. Canvas LMS architecture overview [Electronic resource]. – Access mode: <https://www.instructure.com/canvas> (date of application: 08/23/2025).
9. Blackboard Inc. Blackboard Learn SaaS: Technical Architecture [Electronic resource]. – 2021. – Access mode: <https://help.blackboard.com/Learn/Administrator/SaaS/Deployment> (date of application: 08/23/2025).
10. Open edX. Deploying Open edX using Tutor and Kubernetes [Electronic resource]. – 2023. – Access mode: <https://docs.tutor.overhang.io/> (date of application: 08/23/2025).
11. Bernstein D., Ludvigson E., Sankar K., Diamond S., Morrow M. Blueprint for the Intercloud – Protocols and Formats for Cloud Computing Interoperability // Proc. of the 2009 4th Int. Conf. on Internet and Web Applications and Services. – 2009. – P. 328–336. – DOI: 10.1109/ICIW.2009.55
12. Zaharia M., Chowdhury M., Franklin MJ, Shenker S., Stoica I. Spark: Cluster computing with working sets // Proc. of the 2nd USENIX Conf. on Hot Topics in Cloud Computing. – 2010.
13. Mell P., Grance T. The NIST definition of cloud computing. – National Institute of Standards and Technology, US Department of Commerce. – 2011. – [Electronic resource]. – Access mode: <https://nvlpubs.nist.gov/nistpubs/Legacy/SP/nistspecialpublication800-145.pdf> (date of application: 08/23/2025).
14. European Union. General Data Protection Regulation (GDPR) // Official Journal of the European Union. – 2016. – L119.
15. EDUCAUSE. 7 Things You Should Know About SaaS LMS [Electronic resource]. – 2020. – Access mode: <https://library.educause.edu/> (date of application: 08/23/2025).

METHODS OF INTRODUCING DIGITAL TRANSFORMATION INTO THE ORGANIZATIONS DEVELOPMENT

Molokanova V., Kozyr S.

The processes of development of organizations are considered in two aspects: information (process) and human-centred (behavioural). The information approach involves codification of facts based on certain procedures using appropriate technologies of exchange, distribution and reuse of data. The behavioural approach states that actions of users to place information according to certain rules allow them to successfully transform accumulated knowledge and experience into company profit. The article shows that the introduction of new technologies should be considered from the standpoint of a fundamentally new paradigm – as a convergent science of organizations development.

1. Problem formulation

In today's world digitalization has affected all business sectors, as well as the methodology of managing organizations. Since technology has provided humanity with an unprecedented level of communication between itself and the whole world, the implementation of new business models in practice changes all components of economic activity, which makes it possible to build local, national and global economies based on new values. Thus, the development of new digital tools increasingly requires managers for innovative modelling technologies, big data analysis, the proliferation of social networks and educational platforms – all this opens up unlimited possibilities for virtual management and communication with customers. Now project managers can quickly present their products and services to all stakeholders at the touch of a key. In addition, through the use of comprehensive cloud storage, stakeholders can be quickly integrated into the decision-making process. All this allows us to quickly calculate potential changes and, thus, save resources. At the same time, the new digital approach in no way rejects the achievements of classical management, but only helps to understand its deep aspects.

2. Analysis of recent researches and publications

With the current rapid changes in the environment, modern enterprises face an acute problem of improving their own management system [1]. For a long time, functional and process approaches were used to analyse the development of systems [2, 3]. However, over time, experts have come to the conclusion that the

entrepreneurs who pay less attention to financial indicators and concentrate more on creating organizational value get better results [4, 5].

Burdened by the legacy of old technologies and bureaucratic constraints, incumbents are already facing serious challenges on the path of digital transformation [6]. We expect corporate entrepreneurship to add a more holistic view of internal aspects during the digital transformation process, such as enhancing the impact of knowledge and organizational learning.

Technology as a major determinant of organizational form and structure has been well recognized by scholars for a long time after a significant decline in interest in these relationships until the mid-1990s [7]. Digital technologies are considered the main asset for organizational transformations, given their disruptive nature and systemic effects [8]. To achieve successful digital transformation, changes must occur at different levels of the organization, including adaptation of the core business [9], sharing resources and opportunities [10], intertwining processes and structures [11], adjustments in leadership [12].

Thus, digital transformation is a reality that is becoming increasingly complex due to the special place occupied by knowledge, science, technology and innovation. Both in the field of production and in the service sector, business development increasingly depends on the knowledge generation, qualified information, skills and integral competencies of managers. Here, integral competence is defined as a convergent (consolidating) phenomenon and a significant component of modern economic and technical power, material well-being and sustainable development.

In the process of European integration of various systems into a single global space, a convergent analysis of these systems is very important. It is necessary to look for common ground in perception and approaches, with different ways of thinking, goals, values, directions for further development of the personality and educational systems in which it is formed [13]. Convergence in education is a stabilizing factor that contributes to the selection of optimal components, tested in the process of long-term practice. A critical approach to convergence is characteristic of all countries, forms the basic requirements for science – the preservation and enrichment of human values.

In recent years, there has been intense debate about the importance of knowledge management in our society. Researchers and observers in the field of economics and management science agree that there has been an information transformation and «knowledge» has come into the spotlight [14]. As digital transformation involves the integration of digital technologies into all aspects of business operations to improve products and services, and introduce new revenue

streams, digital transformation has become a major strategy for organizations seeking to reimagine their operations and business models for the modern business era [15-16]. Digital transformation is an organization-wide strategy that aims to use digital technology to modernize key business processes and introduce new services that better engage customers, support employees, improve operations, and add business value. The aspect of strategic renewal in corporate entrepreneurship has different names such as strategic changes, revival, project transformation, reorganization or organizational renewal, which provides a promising direction for digital transformation [17–18].

The article is aimed at studying the manifestations of technological breakthroughs in digital transformation and their impact on the development of technology implementation methodology and the creation of new knowledge on the ways of organizations development.

3. Main part

The theory of project management and the organizations development are closely related, as they are used in the transformation process of the system's transition from one state to another. This process can be considered in two aspects: informational (process) and human-centred (behavioural). The information (process) approach involves the codification of facts on the basis of certain procedures using appropriate technologies of exchange, distribution and reuse of data. The weakness of this approach is that it is not very effective when it is necessary to obtain new knowledge on the basis of facts. The application of IT technologies to knowledge management directs the actions of users to place information according to certain rules that allow us to successfully find and use it in the future. However, in the absence of social interaction, new knowledge does not arise only due to the processing of known information.

The behavioural approach focuses on building a working environment in the company, in which the processes of knowledge exchange and assimilation are facilitated. Looking at what is happening now in the information environment, everyone understands that people are the bearers of knowledge. It is the value of their accumulated knowledge and experience that ultimately translates into the company's profit. Professional communities act as informal groups of people who meet regularly to exchange ideas and knowledge.

In general, when reviewing and analysing the existing literature, it is possible to divide all articles as having been written in two dimensions: technological and behavioural. Now let's consider digital transformation as human organizational

change and information transformation caused by the implementation of digital technology. The result of such a consideration can be presented as the main topics image of digital transformation research (Fig. 1). Such duality reaffirms the principle that certain objects, systems, or theories can be viewed from two different, often opposing, perspectives that are equal and complementary. This does not imply conflict or contradiction, but rather the ability to see the same phenomenon from different perspectives. Such point of view can be useful for both researchers and practitioners, which will allow a more complete understanding of the features of digital transformation.

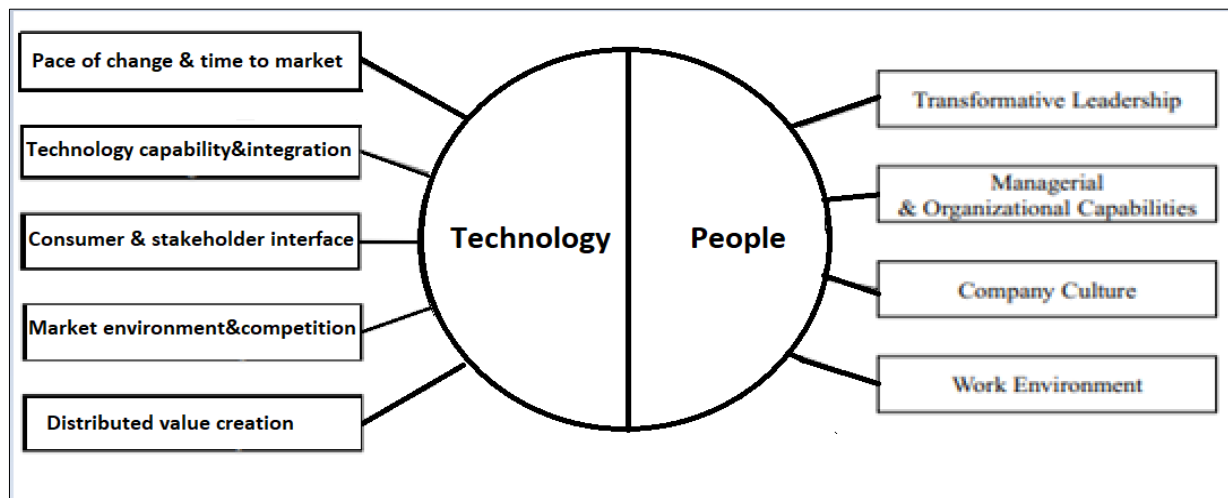


Fig. 1. Two dimensions of Digital transformation

On the technological side the impact of digital technologies is aimed mainly at the consumer interface and market environment. At the same time, understanding the pace of change in times of digital transformation and its direct impact on officials is not yet given enough attention. On the human-centred side researchers mainly focus on leadership and development opportunities in a digital context, while company culture and work environments are still not well understood.

At the macro level, technological aspects are concentrated to digital transformation tools and new business models. Today, innovative ideas can be implemented within a few days, and companies are created literally «overnight». In this sense, in the digital world, the desire to «be the first» and the idea that «winner takes all» has become more important for incumbent firms, since they now have much less time to respond to such threats. Digital companies such as Facebook, Google or Amazon have significantly increased the speed of product launches [19]. With continuous improvements in hardware, software, and communications, such companies are setting the pace for a series of product launches in a short time.

Thus, companies in the digital world are under tremendous pressure to accelerate the adoption of their products as well. In a digitally transformed market, control over the speed of product development and launch is increasingly shifting to an «innovation ecosystem» in the sense of a network of entities with complementary products and services. Many of the authors who represent the technology-focused side emphasize the diffusion of digital technologies as a factor driving business transformation. Such digital technologies can include big data, mobile devices, cloud computing, or search-based applications. Thus, Hess et al. [20] point out that digital transformation «is related to the changes that digital technologies can bring about in the company's business model, which lead to changes in products or organizational structures or to the pace of change and time to market, distributed value, creation and capture of technology».

From other point of view, the human-centred approach emphasizes the role of the individual (manager) in facilitating transformation processes, while facing the problem of balancing at the same time between research and exploitation of resources. At the same time the leaders must believe in the value and benefits of new IT technologies and support their implementation.

We are now seeing some variety of methodologies that digital transformation managers rely on. To bridge the digital divide, various theories are proposed, for example, configuration theory, resource-oriented view, dynamic modelling or business process reengineering. There are other theoretical approaches, such as the view from the perspective of corporate governance and technological breakthroughs through projects. We can see how an exchange of various studies helps to bridge technological gaps. Although we observe that technology-focused research is mostly conceptual, and works focused on human behaviour mainly use case studies. In general, there is little quantitative empirical data in the research literature. This is like a strong indicator for the early stage of digital transformation methodology research.

Historically, digital business transformation began with the implementation of enterprise management information systems, such as enterprise resource planning (ERP) or customer relationship management (CRM). These transformations were usually limited to improving business processes within firms [21]. Today, the possibilities of digital technologies in terms of information processing are incomparably more powerful than in previous time.

Summing up, the current state of digital transformation research can be noted two main events. First, digital transformation restores, blurs, and even dissolves existing industry boundaries, which can lead to cross-industry competition [22].

The old sectoral logic no longer works in times of digital transformation. With the advent of new business models, residents begin to undermine new markets. For example, Google is undermining the mobility sector with its self-driving car subsidiary Waymo, and Amazon has introduced AmazonFresh as a grocery delivery service seen as a potentially tough competitor to supermarkets [22]. Secondly, we see that new sectors of the economy bring new specifics to the market, where old approaches are not applicable, and modern ones become necessary. New creative industries are emerging, characterized by a large number of microbusinesses in the design, cultural, entertainment and media sectors. As different types of innovative networks emerge, new properties of digital infrastructure are needed to support each network. Digital technologies increase the heterogeneity of innovative knowledge networks [23]. The data structure for the technology-oriented dimension is shown in Fig. 2.

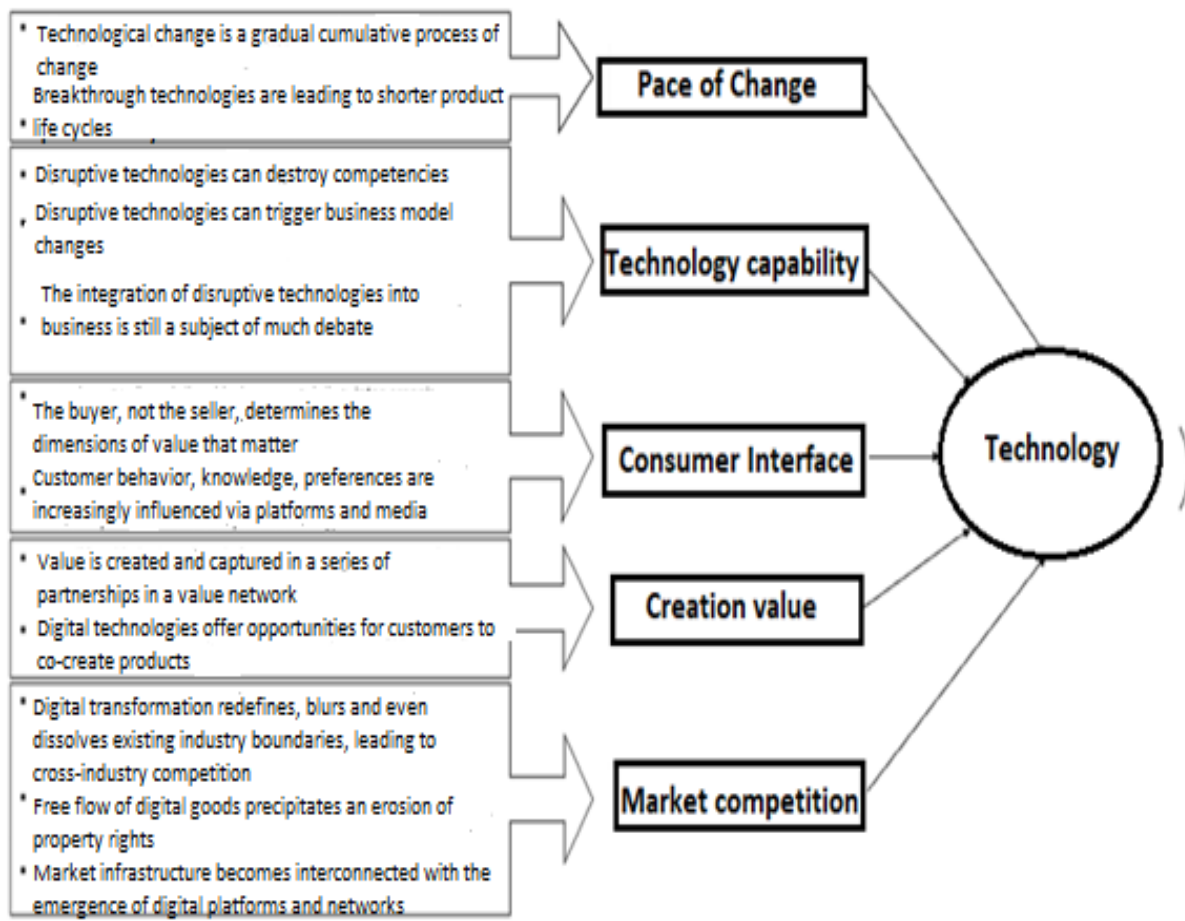


Fig. 2. Data structure for the technology-oriented dimension of disruption

Technologies that drive transformation go far beyond internal process optimization, as they potentially cause drastic changes in business models, organizational strategy, and corporate culture [24]. The role and importance of human

resources are changing, and personal data has become one of the most powerful assets in the digital age [25]. In fact, the impact of the massive increase in the quantity and quality of data generated every day is a game-changer and the possibilities of big data analytics that have yet to be fully felt and understood by society, the economy, and scientists. With regard to the process of dematerialization of tangible products and objects caused by the transformative capabilities of digital technologies, the most remarkable is that in many cases, digital substitutes, such as e-books, demonstrate higher productivity and higher customer benefits than their physical counterparts [26]. Considering the main directions of technology integration most researchers note the importance of new corporate platforms and the scalability of the operational framework as part of a flexible digital infrastructure. The old paradigms of technology integration are no longer effective. It is now necessary to achieve a methodological understanding of how the integration of technologies and transformational measures should be embedded in the organizational structures of existing enterprises. Since it is now the buyer, not the seller, who determines the measurements of the value of goods, companies need to interact with their customers at every stage of the value creation process. In addition, the strong impact of digital technologies on value chains implies a certain degree of deviation from classical analogue business. Now we need new competencies related to the product, platform capabilities, and data value architecture. In addition, modern managers must be ready for new forms of monetization in the digital market. The data structure for the human-centred dimension is shown in Fig. 3.



Fig. 3. Data structure for the people-oriented dimension of technological disruption

4. Conclusions

The focus of digital transformation at the organizational level does not fully reflect the consequences of the digital revolution, as it touches on various aspects of modern life that have not yet been mentioned in our study. Considering the duality of existing research on digital transformation, we find that in the management of organizational development, duality manifests itself as two approaches: technological (informational) and behavioural (human-centred). Their combination opens up significant opportunities for the creation of **a new convergent science**. So, analyzing the directions of future research, we propose to combine information technology and the human approach as two cumulative dimensions of the convergent science of management.

From a macro-level perspective, the literature on digital transformation discusses the topics of changing markets, clusters, and the use of big data. In addition, we note a certain variety of theoretical foundations based on the human-centric approach. We hope that an important contribution of our work is to unite different aspects of digital transformation by integrating knowledge from related disciplinary areas. We hope to continue our research to build a strong methodological foundation for digital business transformation.

References

1. Radulescu, C.Z., Radulescu, M., Filip, F.G., et al. Decision analysis for the project selection problem under risk 9th IFAC Symposium on Large Scale Systems: Theory and Applications 2001.
2. H. Kerzner, Project Management: A Systems Approach to Planning, Scheduling, and Controlling, 11th edition. Hoboken, New Jersey: Wiley, 2013.
3. Delisle C.L. Would the real project management language stand up / C.L. Delisle, D. Olson. Int J Project Manage 22 – 2004 – №4. – pp. 327-337.
4. Garcia A.B. Voting on the agenda: the key to social efficient meetings / A. B. Garcia, J. Kunz, M. Fischer. Int J Project Manage 23 – 2005 – №1. – pp. 17-24.
5. Bushuev, S.D., Bushueva, N.S. «Formation of value in the activities of project-oriented organizations.» Coll. Sciences. Works. Ed. V. A. Racha. – 2009 – №3 (31). – P. 5–14.
6. Lavie D (2006) Capability reconfiguration: an analysis of incumbent responses to technological change. Acad Manag Rev 31(1):153–174
7. Zammuto RF, Griffith TL, Majchrzak A, Dougherty DJ, Faraj S (2007) Information technology and the changing fabric of organization. Org Sci 18(5):749–762
8. Besson P, Rowe F (2012) Strategizing information systems-enabled organizational transformation: a transdisciplinary review and new directions. J Strateg Inf Syst 21:103–124P.
9. Nambisan S, Lyytinen K, Majchrzak A, Song M (2017) Digital innovation management: reinventing innovation management research in a digital world. MIS Q 41(1):223–238. DOI: <https://doi.org/10.25300/MISQ/2017/41:1.03>

10. Besson P, Rowe F (2012) Strategizing information systems-enabled organizational transformation: a transdisciplinary review and new directions. *J Strateg Inf Syst* 21:103–124.
11. Resca A, Za S, Spagnoletti P (2013) Digital platforms as sources for organizational and strategic transformation: a case study of the Midblue project. *Theor Appl Electron Commer Res* 8(2):7.
12. Hansen R, Sia SK (2015) Hummel's digital transformation toward omnichannel retailing: key lessons learned. *MIS Q Exec* 14(2):51–66
13. Managing the level of divergence and convergence in education. *Economics and organization of management*. 2018. № 2. Pp. 14–21. (in Ukrainian).
14. Mårtensson M. A critical review of knowledge management as a management tool. *Journal of Knowledge Management*. 2010. Vol. 4. Iss. 3. P. 204–216. DOI: 10.1108/13673270010350002
15. George B. VUCA 2.0. A strategy for steady leadership in an unsteady world. *Forbes*. 2017. Feb. 17. URL : [https:// www.forbes.com/sites/hbs](https://www.forbes.com/sites/hbs)
16. Dess GG, Ireland RD, Zahra SA, Floyd SW, Janney JJ, Lane PJ (2003) Emerging issues in corporate entrepreneurship. *J Manag Stud* 29:351–378.
17. Dijkman RM, Sprenkels B, Peeters T, Janssen A (2015) Business models for the Internet of Things. *Int J Inf Manag* 35(6):672–678.
18. Zahra SA (2015) Corporate entrepreneurship as knowledge creation and conversion: the role of entrepreneurial hubs. *Small Bus Econ* 44(4):727–735.
19. Bharadwaj A, El Sawy O, Pavlou P, Venkatraman N (2013) Digital business strategy: toward a next generation of insights. *MIS Q* 37:471–482.
20. Hess T, Matt C, Benlian A, Wiesböck F (2016) Options for formulating a digital transformation strategy. *MIS Q Exec* 15(2):123–139.
21. Kaufman RJ, Walden EA (2001) Economics and electronic commerce: survey and directions for research. *Int J Electric Commun* 5(4):5–116.
22. Weill P, Woerner SL (2015) Thriving in an increasingly digital ecosystem. *MIT Sloan Manag Rev* 56(4):27–34.
23. Lyytinen K, Yoo Y, Boland RJ Jr (2016) Digital product innovation within four classes of innovation networks. *Inf Sys J* 26(1):47–75.
24. El Sawy OA, Kræmmergaard P, Amsinck H, Vinther AL (2016) How LEGO built the foundations and enterprise capabilities for digital leadership. *MIS Q Exec* 15(2):141–166.
25. Ng IC, Wakenshaw SY (2017) The Internet-of-Things: review and research directions. *Int J Res Market* 34(1):3–21.
26. Loebbecke C, Picot A (2015) Reflections on societal and business model transformation arising from digitization and big data analytics: a research agenda. *Strategy Inf Syst* 24(3):149–157.

INFORMATION TECHNOLOGY FOR ESTIMATING THE SIZE OF SOFTWARE PROJECTS

Oriekhov O., Farionova T.

The aim of the study is an information technology development for automated processing of the information from source code metrics for the software size estimation using nonlinear regression models. The object of the study is the process of the source code metrics information processing for the software size estimation using information technology. The subject of the study is information technology for source code metrics information processing for the software size estimation. A review of the usage of software development effort estimation displays that nowadays, software development companies and teams start to use effort estimation in their CI/CD tools and require efforts from DevOps engineers to integrate the software solutions in their IT infrastructure. The study offers a software-as-a-service information technology solution for software size estimation for integration into the CI/CD infrastructure of software development companies. A system design is given as a sequence diagram and a class diagram. The information technology is implemented as a software-as-a-service solution to offer possibilities for integration into the CI/CD infrastructure of software development companies.

Introduction

Estimating the size of a software project is a critical task in software development effort estimation. Accurate size estimation helps project managers plan resources, set deadlines, and control budgets. Traditionally, software size estimation in thousands of lines of code (KLOC) is used in parametric models such as COCOMO, COCOMO II, COSYSMO, SLIM, etc. Software development effort estimation in a KLOC basis is more accurate in comparison with functional points, because the KLOC variant takes into account such environmental factors as programming language and software category and offers alignment of the efforts with the evolution of the programming language. A disadvantage of the KLOC usage is the requirement to have relevant mathematical models for an accurate estimation, because some math models could be missing or to be obsolete with intensive evolution of the programming language [1].

Software development effort estimation remains an actual task that every software development company faces before starting any new software development project or continuing to implement existing software products. The Standish Group report [2] indicates a moderate positive trend in the share of successfully completed projects. In 1994, only 16% of software projects were successfully implemented, meeting deadlines, staying in the budget range, and fulfilling all declared

requirements. By 2020, the share of successfully implemented projects had risen to 35%. The percentage of failed projects decreased from 31% in 1994 to 19% in 2020. The share of projects which are facing challenges, such as exceeding deadlines, overusing budgets, or failing to meet all declared functional requirements, has shown a downward trend of approximately 4%. Also, the research [2] demonstrates a strong correlation between project size and failure probability or potential challenges and risks during software project development. Larger software projects are more difficult to manage because they involve bigger resource allocation, wider coordination, and they have higher probability of unforeseen obstacles. The CHAOS Report findings show that large projects require more accurate estimation models and risk management strategies to improve their chances of success. Without accurate effort and size estimation at the initial planning stage, the expectation may be mismatched with actual development realities and, it leads to an increased probability of project failure.

Nowadays, traditional software exists for software development effort estimation: SEER-SEM, COCOMO II, QSM SLIM, EstimateTool, IBM Rational Tools, etc and, alongside with traditional software, there are continuous integration and continuous delivery (CI/CD) tools for software size and software development effort estimation are demanded. However, as software development becomes more dynamic and continuous, especially with the use of Agile and DevOps practices, there is a growing demand for more flexible and integrated software estimation tools. In this context, CI/CD tools are being used not only to automate building, testing, and deployment, but also to support software size and software development effort estimation in real-time. Integrating SDEE tools into CI/CD pipelines allows teams to automatically estimate effort when changes are made to a UML class diagram, a codebase is modified, a new feature is added, or the existing code is refactored. These estimates can be generated and updated continuously throughout the development process. Moreover, the integration enables comparison between expected and actual development effort in an automated manner, supporting better planning, prediction, and project control without manual input [3-5].

As we can see, active growth and injection of CI/CD tools in the software planning and development process requires transformation of the traditional SDEE tools to offer services that could be injected in automation of the software size and development effort estimation process at the IT infrastructure side. To achieve the automated code metrics processing and usage of modern tools for the software size and effort estimation, it is necessary to propose scalable, easy to update and maintain software solutions. In this case, to provide appropriate and relevant math models,

the Software-as-a-Service (SaaS) distribution model guarantees that modern software size estimation tools will always be up to date, can be easily integrated into the IT infrastructure, and ready for use in CI/CD pipelines for software development efforts assessment.

Therefore, the aim of the study is an information technology development for automated processing of the information from source code metrics for software size estimation using nonlinear regression models. The object of the study is the process of the source code metrics information processing for the software size estimation using information technology. The subject of the study is information technology for source code metrics information processing for the software size estimation.

Modeling of the information technology for software size estimation

The SaaS distribution model was chosen to provide easy access to software size calculation possibilities. It ensures fast updates, maintenance, and improvement of the software application without user intervention, providing instant access to new mathematical models, software features, security enhancements, and bug fixes; providing an application programming interface (API) for integrating KLOC estimation into the software tools of customer's services or existing IT infrastructure, allowing for the creation of comprehensive solutions combined into a single infrastructure. This approach ensures high compatibility and flexibility, allowing organizations and companies to use calculation tools as part of their IT environment, which improves the efficiency of analytical assessment of software development projects [6].

The SaaS information technology for processing information from code metrics to evaluate the size parameter of the software project should solve the following tasks:

- 1) information code metrics processing for estimating the average KLOC software size of software applications, depending on a specific programming language in the early stages of software project planning, using mathematical models that are based on quantitative code metrics.
- 2) the confidence interval bounds estimation of the average KLOC size of software applications depending on a given confidence probability;
- 3) the prediction interval bounds estimation of the average KLOC size of software applications depending on a given confidence probability;
- 4) providing an API for integrating software into customer systems;

5) returning a list of existing mathematical models for estimating information from software code metrics depending on a programming language, that declares necessary code metrics and their acceptable value ranges.

Mathematical background for the IT for estimating the size of software applications considers mathematical models for KLOC estimation depending on a programming language and independent factors that could be obtained from UML class diagrams. The three-factor nonlinear regression models for Data Science and Machine Learning JAVA applications size estimation, the five-factor non-linear regression models for JAVA applications size estimating, and three-factor non-linear regression model for Kotlin applications size estimation are selected for the initial integration into the SaaS information technology [7–9]. These models were built upon modern open-source software projects and offer software size estimation for the highly used JAVA programming language, and intensively growing Kotlin programming language.

The target audience of the SaaS is software development engineers, DevOps, scientists, and project managers. The users have the necessary experience to obtain software size estimates and integrate it into the IT infrastructure of a company using the SaaS variant and the declared documentation. The entire audience is declared as a user role for the IT. Another role is an administrator to provide access to manage the IT, and it will be omitted in the further use case description below.

The user role has the following specification of **use cases** of the SaaS:

1) Obtain information about mathematical models

The user receives a list of available mathematical models that can be used for calculations or obtains information about a mathematical model by its identifier. The list of parameters includes the programming language, the category of the mathematical model, or the identifier of the mathematical model in the system.

The main flow of the use case:

- the user sends a request for information about mathematical models;
- the system checks the request parameters and proceeds with it;
- the system forms a list of available mathematical models in JSON format or a single model in JSON format if the request was made using a unique model identifier.

Alternative flows:

- if the request or input data is incorrect, the system returns an error (HTTP 400 Bad Request);

- if the mathematical model is missing, the system returns an error (HTTP 404 Not Found);
- if the user is not authorized, the system returns (HTTP 401 Unauthorized).

2) Calculate the software size KLOC value using the selected mathematical model

The user forms a request with the values of the declared independent factors of the mathematical model and the significance levels α for calculating prediction intervals and confidence intervals. The system returns the calculation results. The pre-condition of the use case – the math model is selected for the calculation.

The main flow of the use case:

- the user selects the mathematical model's identifier for the query;
- the user sends the input data form for calculating the dependent factor;
- the system validates the query parameters and performs calculations of the dependent factor, as well as prediction intervals and confidence intervals;
- the system forms the calculation result in JSON format.

The SaaS platform's estimation workflow is presented in Figure 1 in the form of a UML activity diagram that illustrates behaviour of the two use cases. It visualizes the sequential interaction between use case № 1, «Obtain information about mathematical models», and use case № 2, «Calculate the software size KLOC value using the selected mathematical model». The diagram highlights the flow of actions and decision points required to retrieve model parameters and perform calculation of software size estimates.

The declared workflow of the SaaS offers all necessary opportunities to integrate the IT solution into CI/CD infrastructure of the software companies and build continuous software size and software development effort estimation solutions.

Overview of the SaaS system design

Information support for the SaaS solution for code metrics information processing includes a software development solution, a database, and predefined data models' structures designed to store the parameters of nonlinear regression models for evaluating the size of software applications and the parameters for evaluating the bounds of confidence and prediction intervals. This approach allows adding new models or refining existing parameters of the math models without recompiling the software solution, which allows avoiding the CI/CD flow breaking.

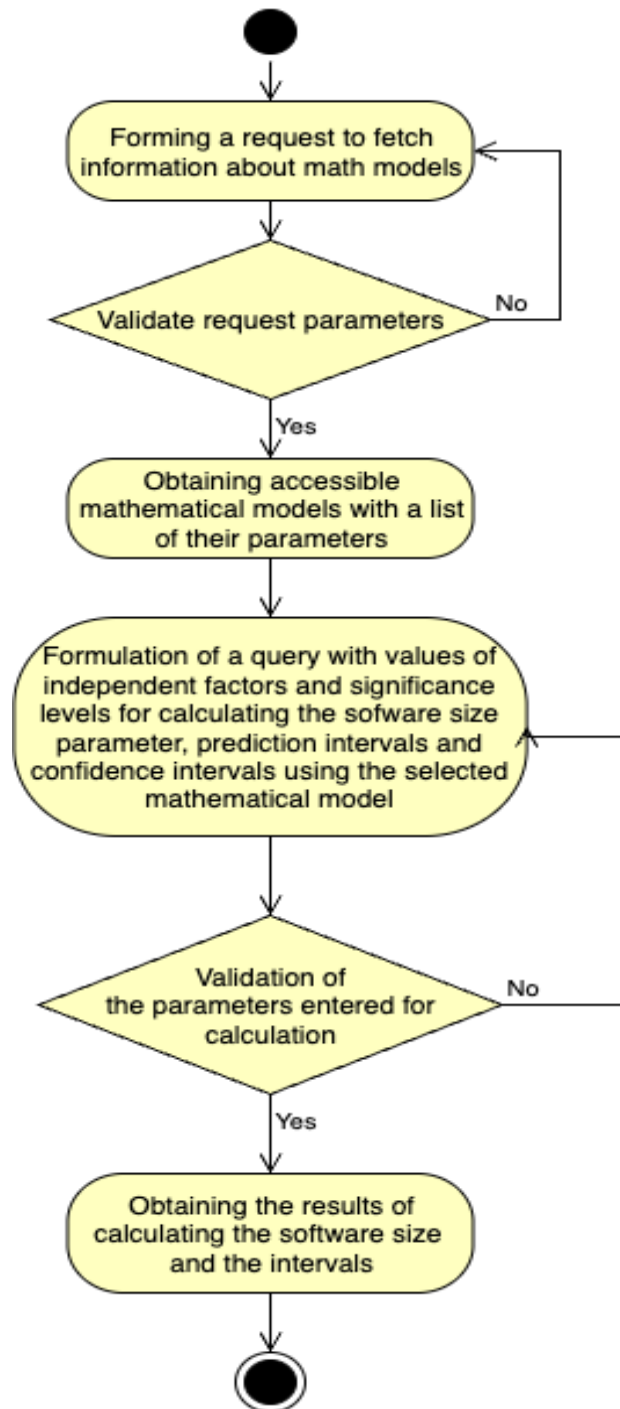


Fig. 1. UML-activity diagram of the software size estimation workflow

The Kotlin programming language with the Spring Framework was chosen as the programming language for developing the estimation service to build network layer interaction. Kotlin is a high-level object-oriented programming language that is distinguished by its speed and performance, and it is fully compatible with the JAVA programming language, which is a great advantage, since a large number of libraries and frameworks have been developed in JAVA, allowing for the implementation of a variety of functionalities, including mathematical calculations [10].

Spring Framework was chosen as the framework for developing the web application due to its advantages, such as a free distribution model, a wide range of tools for implementing network functionality and support for REST API, quick startup and automatic configuration, which reduces the amount of manual configuration for various network interaction components, integration with databases, provides support for working with JSON data, resource access control and support for standards such as OAuth2 and JWT (JSON Web Tokens), which is especially important for network applications such as SaaS, where data security is critical [11].

Technical support for the SaaS information technology has the following hardware requirements: CPU with at least 1 GHz, recommended CPU should have 2 GHz or higher, RAM size should be minimum 512 MB, recommended 1 GB or higher, minimum disk space is 500 MB of free space, internet connection. The hardware requirements are based on the JVM requirements [12].

The user's hardware environment requires uninterrupted access to the Internet. The user's hardware and software must allow requests to be made via the HTTPS protocol (SSL/TLS support) over the Internet [13].

The logical model of the IT takes into account different programming languages, categories, and multivariate parameters structure of the math models. It includes a unique identifier to store all different models in the system, programming language, category, math model name, math model description, quantity of the independent factors, list of the independent factors, dependency factor, datetime and reference when the model were published, accuracy testing information, model restriction information, and a complex data structure to store the math model coefficients and constants.

A UML class diagram was created for the SaaS. The class diagram shows the main entities of the system, their relationships, and interactions between them. The class diagram is shown in Figure 2.

The software information processing SaaS has the following requirements for the software on the server-side environment [12].

Operating system requirements for running the SaaS application:

- Windows: Windows 10, Windows 11, Windows Server 2016, 2019, 2022;
- Distributions that support glibc 2.17+ (e.g., Ubuntu 18.04+, CentOS 7+, etc.);
- Mac OS: macOS 10.15 (Catalina) and newer.

Runtime environment requirements are Java Runtime Environment 8 or higher, PostgreSQL database server version 14 or higher, and OpenSSL version 1.1.1 or newer.

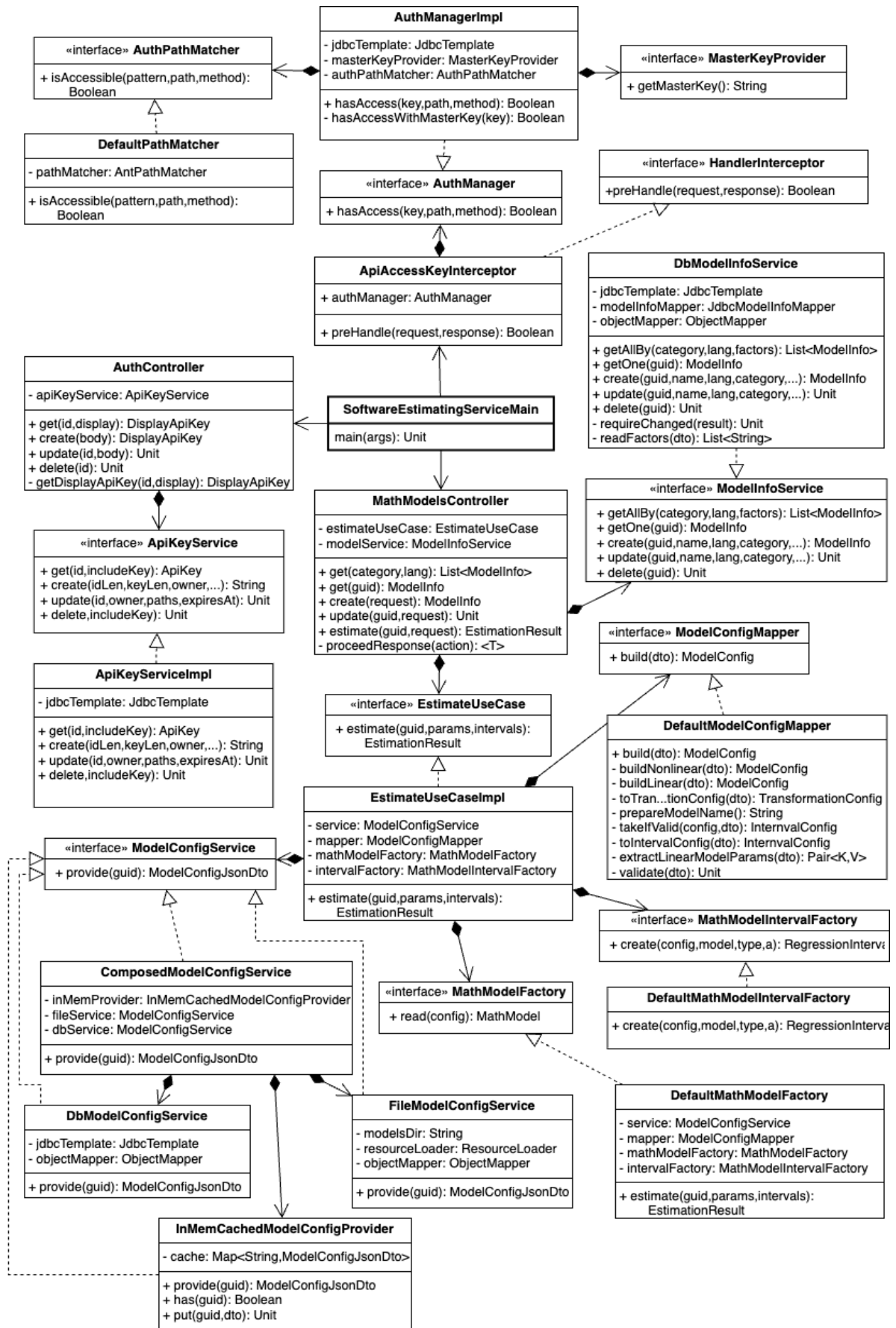


Fig. 2. Class diagram of a SaaS application for software size estimation from code metrics

Conclusion

The results of the study propose an automated processing of the information from code metrics for software size estimation at an early stage of the software project development and provide software tools solutions for integrating the calculator in CI/CD infrastructure of a company for software size estimation or effort monitoring systems.

The practical significance of the obtained results allows us to recommend the described software construction, or the software estimation tools for use in practice for integration in software development life cycle monitoring tools for use with COCOMO II, COSYSMO, etc, to obtain effort estimates.

Prospects for further research may include extending the software metrics estimation service software with a wider range of models for estimating the size of projects for different programming languages and adding software development effort estimation models like COCOMO and COCOMO II.

References

1. Farionova, T. A., & Oriekhov, O. S. (2024). A review of software development effort estimation methods. *Collection of Scientific Papers of Admiral Makarov National University of Shipbuilding*, 494, 102–111. DOI: [https://doi.org/10.15589/znpp2024.1\(494\).15](https://doi.org/10.15589/znpp2024.1(494).15)
2. Johnson, J., & Mulder, H. (2021). Endless Modernization: How Infinite Flow Keeps Software Fresh. *The Standish Group, Hans Mulder's Lab*. URL: https://www.researchgate.net/publication/348849361_Endless_Modernization_How_Infinite_Flow_Keeps_Software_Fresh
3. Meedeniya, D., & Thennakoon, H. (2021, August). Impact factors and best practices to improve effort estimation strategies and practices in DevOps. In *Proceedings of the 11th International Conference on Information Communication Management (ICICM 2021)*, pp. 11–17. ACM. DOI: <https://doi.org/10.1145/3484399.3484401>
4. Tyagi, A. (2021). Intelligent DevOps: Harnessing Artificial Intelligence to Revolutionize CI/CD Pipelines and Optimize Software Delivery Lifecycles. *Journal of Emerging Technologies and Innovative Research*, 8, 367–385.
5. Port, D., Taber, B., & Emkani, P. (2024). Investigating effectiveness and compliance to DevOps policies and practices for managing productivity and quality variability. *Journal of Systems and Software*, 213. DOI: <https://doi.org/10.1016/j.jss.2024.112030>.
6. Golding, T. (2024). *Building Multi-Tenant SaaS Architectures*. O'Reilly Media.
7. Oriekhov, O., & Farionova, T. (2024). Nonlinear regression models for software size estimation of Data Science and Machine Learning JAVA-applications. In *Proceedings of the V International Workshop IT Project Management (ITPM 2024)*, Bratislava. URL: http://dx.doi.org/10.23939/IW_itpm2024.054
8. Oriekhov, O., & Farionova, T. (2025). Five-factor nonlinear regression model for JAVA software size estimation. In *Proceedings of the V International Workshop IT Project Management (ITPM 2025)*, Kyiv.

9. Prykhodko, S. B., Prykhodko, N. V., & Koltsov, A. V. (2024). A nonlinear regression model for early LOC estimation of open-source Kotlin-based applications. *Radio Electronics, Computer Science, Control*, 85(1), 85–95. DOI: <https://doi.org/10.15588/1607-3274-2024-1-8>.
10. JetBrains. (2025). *Kotlin programming language*. KotlinLang. URL: <https://kotlinlang.org/>
11. Spring. (n.d.). Spring Boot. URL: <https://spring.io/projects/spring-boot>
12. Oracle. (n.d.). *Oracle JDK 8 and JRE 8 Certified System Configurations*. URL: <https://www.oracle.com/java/technologies/javase/products-doc-jdk8-jre8-certconfig.html>.
13. Rescorla, E. (2018). *The transport layer security (TLS) protocol version 1.3 (RFC 8446)*. Internet Engineering Task Force. DOI: <https://doi.org/10.17487/RFC8446>

**MATHEMATICAL BASES OF THE METHODOLOGY
FOR BUILDING AN INTELLECTUAL SYSTEM BASED ON THE THEORY
OF INTELLIGENCE AND THE APPARATUS OF PREDICATE ALGEBRA**

Plehova G., Sharonova N., Neronov S., Plehov D., Fedorovych V.

Introduction

A method of comparative identification and tools for building an intelligent system are proposed. It is shown that finite predicate algebra (FPA) is a universal mathematical apparatus for the formal description of deterministic, discrete, finite objects of various nature. FPA makes it possible to interpret knowledge in a rigorous mathematical form, where different features and their meanings are related to each other by means of Boolean and predicate operations. It is shown that when constructing a feature system for a knowledge base, it is advisable to use an algebraic-logical model of the subject area. It is shown that in practical problems related to the interpretation of knowledge, which is mostly presented in textual form, it is not necessary to obtain all sets of values of semantic features, but it is necessary to obtain one or several essential sets of feature values (target variables), which are a fixed set of values of other features. When solving such problems, other variables that are not included in the initial conditions and are not target variables are excluded from the equation by linking them to existential quantifiers. An algorithm for excluding non-essential variables when solving algebraic-logical equations is proposed, which reduces the number of necessary calculations. Finite predicate equations of various types are considered. A description of classification problems based on predicate equations is given. The problem of classifying objects based on features that take discrete values is described mathematically as a solution of predicate equations. A broad class of predicates is described from which redundant variables can be removed and focus can be placed on the relationships between essential variables. A method for removing non-essential variables by applying an existential quantifier is proposed and demonstrated using a real-world example. An example is given of the use of the mathematical apparatus of intelligence theory, the method of comparator identification, and the tools of finite predicate algebra, which consists in constructing a model for the identification of medical diagnostic parameters in the form of a system of predicate equations, the solution of which provides an interpretation of medical knowledge in a given field.

1. Comparator identification – a method of intelligence theory

The need to consider the theory of intelligence as a *formal doctrine* is due to the fact that it requires a special mathematical language that is not sufficiently developed in existing branches of mathematics. Therefore, the theory of intelligence, along with a substantive study of the human mind, is also forced to develop the necessary formal apparatus. The possibility of teaching the theory of intelligence deductively, based solely on physically observable facts, is based on the method of axiomatic description of the human mind developed by the scientific school of Prof. Shabanov-Kushnarenko Yu.P. and his followers. This is a method of comparison or *comparator identification method*. The essence of the method lies in the fact that the test subject (the person whose intellect is being studied) forms the meaning of certain predicates through their physical reactions in specially designed experiments $P_1, P_2, \dots, P_r$.

These experiments reveal the properties of predicates $P_1, P_2, \dots, P_r$, which are formally written as logical equations linking predicate variables $X_1, X_2, \dots, X_r$. Some of these equations are used as axioms or initial postulates of intelligence theory. From the *axioms*, as from equations, the values of the predicate variables $X_1, X_2, \dots, X_r$ are found, which are, respectively, the predicates $P_1, P_2, \dots, P_r$. The internal structure of the found predicates characterizes certain details of the mechanism of intelligence.

The comparison method was first used by Newton in his physical study of human color vision. Acting as a test subject, he observed arbitrary light emissions X_1 and X_2 in the fields of comparison and recorded the equality or inequality of their colors. The predicate $P(X_1, X_2)$ formed in this way was first linked by logical equations-axioms by Grassmann. Based on Grassmann's postulates (laws), Schrödinger was the first to construct a deductive theory of human color vision.

When studying human intelligence using the comparison method, the researcher influences the subject's sensory organs with physical signals (stimuli) $X_1, X_2, \dots, X_n$, which generate certain subjective experiences (states) $Y_1, Y_2, \dots, Y_n$ in the subject's consciousness. It is assumed that states $Y_1, Y_2, \dots, Y_n$ depend unambiguously on the corresponding stimuli $X_1, X_2, \dots, X_n$. This means that there are (exist) functions

$$Y_1 = f_1(X_1), Y_2 = f_2(X_2), \dots, Y_n = f_n(X_n).$$

In experiments, the stimuli $X_1, X_2, \dots, X_n$ are taken from sets $A_1, A_2, \dots, A_n$ clearly defined by the researcher, so that always $X_1 \in A_1, X_2 \in A_2, \dots, X_n \in A_n$.

The researcher selects the sets A_i ($i \in \{1, 2, \dots, n\}$) arbitrarily, at his discretion, based on the scientific tasks he sets for himself. It is assumed that each of the stimuli $X_i \in A_i$ generates a specific state Y_i . The set of all values of the function $Y_i = f_i(X_i)$, defined on the set A_i , is denoted by the symbol B_i . Thus, each of the functions f_i is a surjection mapping the set A_i onto the set B_i . The functions f_i characterize the ability of the test subject to respond to external objects with corresponding subjective states.

2. Finite predicate algebra as a universal means of formal description for objects of different nature

It is believed that finite predicate algebra is a universal mathematical apparatus for the formal description of deterministic, discrete, finite objects [1–11]. In works [4, 5], it was proven that in order to develop the theory of intelligence, it is first necessary to have a formal language that can be used to mathematically describe the structures and functions of human intelligence, especially those related to the accumulation and interpretation of knowledge by humans. The scientific school of Professor Yu.P. Shabanov-Kushnarenko and his followers uses *finite predicate algebra* (FPA) as such a language.

It should be noted that *deterministic processes* are processes with a unique result, in which there is no element of chance. *Discrete processes* are processes in which information takes the form of separate portions or quanta – digits, letters, words, formulas, etc. – in which there is no factor of continuity. *Finite processes* are processes in which only a finite number of units of information are involved, in which there is no factor of infinity. By definition, the theory of intelligence must be limited to the study of purely machine-like functions of human intelligence, which are characterized by *determinism, discreteness, and finiteness*. This raises questions about the extent to which the imposed restrictions narrow the class of human intelligence functions that can be modeled within the framework of the theory of intelligence.

In order to be able to describe the functions of intelligence mathematically, it was necessary to create a formal language in which such a description could be made. The formal language had to be chosen so that any finite alphabet operator could be written in a convenient form. Such a language is the FPA described below [6–8].

The concept of a *finite predicate* was introduced as follows. Let A be a finite alphabet consisting of k letters $a_1, a_2, \dots, a_k$, Σ be a set consisting of two elements denoted by the symbols 0 and 1 and called *false* and *true*, respectively. A variable

defined on the set A is called a (*letter*) *literal*, and a variable defined on the set Σ is called a *logical* variable.

A *finite n -place predicate* over an alphabet A is any function $f(x_1, x_2, \dots, x_n) = t$ of n letter arguments $x_1, x_2, \dots, x_n$, defined on the set A , and which takes the logical values t . Sometimes a finite predicate f is called *k -ary*, emphasizing that its alphabet A consists of k letters.

Each finite alphabet operator can be assigned a specific finite predicate. This can be done as follows. Let

$$F(x_1 x_2 \dots x_m) = y_1 y_2 \dots y_m$$

– is an arbitrarily chosen finite alphabet operator that converts input words $x_1 x_2 \dots x_m$ of length m into output words $y_1 y_2 \dots y_m$ of the same length, consisting of letters of the alphabet A . Let us construct a $2m$ -place finite predicate $f(x_1, x_2, \dots, x_m, y_1, y_2, \dots, y_m)$ over the alphabet A , using the following rule:

$$f(x_1, x_2, \dots, x_m, y_1, y_2, \dots, y_m) = \begin{cases} 1, & \text{if } \Leftrightarrow F(x_1 x_2 \dots x_m) = y_1 y_2 \dots y_m \\ 0, & \text{if } F(x_1 x_2 \dots x_m) \neq y_1 y_2 \dots y_m. \end{cases}$$

Let's write down the equation:

$$f(x_1, x_2, \dots, x_m, y_1, y_2, \dots, y_m) = 1.$$

This equation links the variables $x_1, x_2, \dots, x_m, y_1, y_2, \dots, y_m$ by a certain relation. Substituting into $f(x_1, x_2, \dots, x_m, y_1, y_2, \dots, y_m) = 1$, the letters $x_1, x_2, \dots, x_m$ of the input word $x_1 x_2 \dots x_m$ of the alphabet operator F , we obtain as a result of solving this equation the letters $y_1, y_2, \dots, y_m$ of the output word $y_1 y_2 \dots y_m$. Thus, the predicate f , constructed in this way, contains all the information about the alphabet operator F that interests us. With its help, we can determine the output word of the alphabet operator represented by this predicate for any input word.

In order to be able to write finite predicates in the form of formulas, a special algebraic system called FPA is introduced. More precisely, not one algebra is introduced, but a whole family of such algebras. *The type of finite predicate algebra* is characterized by an *alphabet of letters* A consisting of symbols $a_1, a_2, \dots, a_k$, and an *alphabet of variables* consisting of n symbols $x_1, x_2, \dots, x_n$. Using the finite predicate algebra with the alphabet $A = \{a_1, a_2, \dots, a_k\}$ and an alphabet of variables $B = \{x_1, x_2, \dots, x_n\}$, any n -place k -ary predicate $f(x_1, x_2, \dots, x_n)$ defined over the alphabet A can be written. In this case, it is assumed that the letters and variables of the alphabets A and B are numbered.

In finite predicate algebra, there is a very special function, or more precisely, a unique predicate called the *recognition predicate*. Each formula of the form $a_i(x_j)$ can be conveniently regarded as a single-place predicate that depends only on the variable x_j and is defined as follows:

$$a_i(x_j) = \begin{cases} 1, & \text{if } x_j = a_i, \\ 0, & \text{if } x_j \neq a_i. \end{cases}$$

The predicate $a_i(x_j)$ «recognizes» a single letter a_i among the possible letters of the alphabet A , and is therefore called the predicate of recognition of the letter a_i or even the recognition of the letter a_i , which depends on the variable x_j . There are k_n different letter recognitions in total. By applying disjunction and conjunction operations to different letter recognitions, repeatedly and in different orders, you can obtain different predicates. Letter recognition is considered to be an *elementary predicate of finite predicate algebra*. The set of elementary operations and elementary predicates forms the basis of *finite predicate algebra*.

It has been proven that the algebra of finite predicates is *complete*. Thus, any finite predicate can be written in the ASP language. The algebra of finite predicates is not only complete, but even *redundant* in a certain sense. The true predicate denoted by formula 1 can be expressed by other means of finite predicate algebra, for example, in the form of a disjunction of all conjunctions of the form $x_1^{\sigma_1} x_2^{\sigma_2} \dots x_n^{\sigma_n}$. In the case when $k \geq 2$, i.e., when the alphabet A contains at least two letters a_1 and a_2 , one can do without formula 0. The predicate denoted by this formula, which is identically false, can be written, for example, in the form of the formula $x_1^{a_1} x_1^{a_2}$.

Most often, *disjunctive algebra* and its various conservative extensions are used as the language for recording finite predicates. The convenience of disjunctive algebra of finite predicates lies in the fact that its language allows for concise and convenient recording of *truth and falsity laws*. These laws play a special role in any finite predicate algebra. In essence, there are requirements that must be satisfied for the correct introduction of variables on finite sets.

The *truth law* specifies the domain of variation of a variable, and the *falsity law* ensures that all elements of the set in which the variable is defined are pairwise distinct. In any other algebra, the *truth* and *falsity laws* would be written in much more cumbersome expressions. Hereinafter, we will refer to the *disjunctive algebra* and its conservative extensions simply as the *finite predicate algebra*.

Thus, in this subsection, we have shown that the finite predicate algebra is a universal mathematical tool for the formal description of deterministic, discrete, finite objects of various kinds.

The convenience of the disjunctive algebra of finite predicates lies in the fact that its language allows for concise and convenient notation of *truth and falsity laws*, which play a special role because they set the necessary and sufficient requirements for the correct introduction of variable attributes on finite sets.

The truth law specifies the domain of variation of a variable, and the falsity law ensures that all elements of the set on which the variable is defined are pairwise different. In any other algebra, the truth and falsity laws would be written as much more cumbersome expressions.

3. Formal description of the subject space in the language of finite predicate algebra

A distinction should be made between the formal description of an object and its abstract description. The abstract description of an object always follows the formal description. This is a stage of more in-depth study of the object compared to its formal description. The formal description of an object is the replacement of the object with mathematical constructions. An abstract description of an object is an axiomatic construction of a theory of those mathematical structures that were obtained in the formal description of the object.

The representation and interpretation of knowledge play an important role in various areas of computer science. Various methods of discrete mathematics are used to formally represent information about objects and processes in knowledge bases. In cases where information about objects and processes in knowledge bases, represented by discrete information features, has a rather complex logical structure, various methods and models of discrete mathematics are used for its formal representation, including logical equations with Boolean variables. Thus, logical methods of image recognition involve the compilation and solution of logical equations with variables that take on values of 1 and 0, depending on whether a given object has a certain property or not. Solving such equations allows either to identify an object by the available sets of attribute variable values or to establish unknown properties of a given object [1].

4. Method for studying the relationships between discrete characteristics of objects

Logical classification methods usually involve compiling and solving logical equations with variables that take on values of 1 and 0, depending on whether a given object has a certain property or not. Solving such equations allows either to identify an object by the available sets of attribute variable values or to establish unknown properties of a given object. A natural generalization of Boolean algebra equations is finite predicate algebra equations, which allow you to operate with arbitrary attribute variables defined on different finite sets. The use of such equations for constructing logical conclusions in knowledge bases allows expanding the capabilities of Boolean logical methods for object recognition and classification. When classifying objects, we deal with sets of features, selecting some of whose values allows us to determine the membership of the object under consideration to a certain class. This section presents a method for investigating the relationships between discrete features of objects. Various types of predicate equations are considered.

When analyzing the relationships between essential features of data, we often encounter rather complex systems of logical equations, which, nevertheless, can be simplified due to their specific properties. To demonstrate the procedure for eliminating non-essential features, a real example from the field of medicine will be considered below.

Many practical problems require the use of logical classification methods. For example, in [16], binary feature vectors are classified. The proposed method can be used in various classification problems in different areas. The classification process boils down to studying logical-dynamic systems depending on some initial states. In [17], correction functions are constructed for logical object recognition methods. An interesting algorithm for logical classification using correction functions is proposed. In [18], logical data classification is used to analyze hyperspectral data. Some combinatorial issues of logical classification are discussed in [19]. The authors applied their research to the classification of proteomic expressions of Alzheimer's disease. In [20], logical algorithmic methods for constructing decision trees are considered. In [21], a logic-based classification method for text recognition is proposed. In [22], a logical approach to machine learning is proposed. The authors developed a special logical classifier. The logical design of a strong classifier using weak classifiers is considered in [23]. A combination of different methods, including logical classification, was tested in [24]. A method for classifying text messages based on logical extraction is considered in [25]. In [26], finite predicate networks for radar detection are investigated. A new approach to logical classification and

recognition is proposed in [27]. Subsystems of Boolean equation systems are investigated in [28]. In [29], a method for distributed solving of logical equations is presented. In [30], some specific types of logical functions are considered. The main focus is on symmetric functions, which are widely used for classification purposes. The issue of entropy in Boolean networks is discussed in [31]. Logic-predicate networks are investigated in [32]. The use of predicates for classifying access problems in the Internet of Things is considered in [33]. Many works devoted to classification often use Boolean algebra, fuzzy logic, or neural networks as basic mathematical tools [19]. Logical approaches to fact-based analysis are discussed in [20].

A common generalization of Boolean algebra equations is finite predicate algebra equations [33], which allow operations on arbitrary attribute variables defined on finite sets. The use of such equations for constructing logical inferences in knowledge bases makes it possible to extend the capabilities of logical methods for object recognition [34]. There are many ways to represent data and the relationships between them. First, we will briefly consider Boolean algebra equations and some problems that arise when solving them.

The axioms of Boolean algebra are satisfied by some mathematical structures that are of great importance for information systems. Such structures include, for example, logic algebra, set algebra, and ASP. Often, the dependencies between elements of Boolean algebra can be described by an equation or a system of equations, where one equation (or a certain number of equations) can contain both unknown elements (variables) and known elements (parameters). It is interesting to find out under what conditions an equation (or a system of equations) has a solution, as well as the conditions under which this solution is unique. Finally, it is necessary to be able to find the solution to the equation.

In addition, if the conditions for uniqueness of the solution are not met, it is sometimes impossible to find all solutions to a given equation or system of equations, and it is practically impossible to try them all in turn, since their number can be very large. However, it would be desirable to be able to study these solutions, select the optimal ones (for specific problems), and consider the structure of these solutions in the same way as the structure of a set of differential equations is considered in the theory of differential equations, despite the fact that this set may be infinite. It is possible to find all solutions to Boolean equations without going through them one by one, but by describing them with formulas so that all solutions to a given equation or system of equations can be obtained from these formulas.

5. Analysis of logical connections between essential features of data in systems of predicate equations

A universal method for solving systems of finite predicate algebra equations is to reduce the predicate, given by the system of equations and initial conditions, to a perfect disjunctive normal form. However, this procedure involves searching through many intermediate solutions, and its practical implementation requires significant computer time. For some types of predicate equations, taking into account the peculiarities of their structure, simpler solution algorithms can be developed.

In many practical problems related to the semantic processing of medical data, natural language information, and customer data, it is not necessary to obtain all sets of values of semantic features, but it is necessary to obtain one or more sets of feature values (target variables) that are of interest to the user. It is often necessary to find the values of target variables under given initial conditions, which are a fixed set of values of other features. When solving such problems, other variables that are not included in the initial conditions and are not target variables are excluded from the equation by binding them with existential quantifiers.

Let us consider the solution of Boolean equations using the traditional method of reduction to PDNF (perfect disjunctive normal form). The essence of this method is that the logic algebra function on the left side of the equation is written in PDNF form, or more precisely, the sets of values of binary variables that satisfy this equation are written directly. For example, let us have the following equation

$$(X_1 \vee \bar{X}_2)(X_2 \vee \bar{X}_3) = 1.$$

Writing down the left side as PDNF, we obtain the equation

$$X_1X_2X_3 \vee X_1X_2\bar{X}_3 \vee X_1\bar{X}_2\bar{X}_3 \vee \bar{X}_1\bar{X}_2\bar{X}_3 = 1,$$

which shows that the solutions are sets $\{1,1,1\}$, $\{1,1,0\}$, $\{1,0,0\}$, $\{0,0,0\}$.

For logic algebra, this approach is universal in a sense, since it is always possible to obtain all solutions to any equation. Suppose that in the above example, X_1, X_2, X_3 – are variable sets of some set algebra, and we interpret conjunction, disjunction, and negation as intersection, union, and addition in set algebra.

For logic algebra, this approach is universal in a sense, since it is always possible to obtain all solutions to any equation. Suppose that in the above example X_1, X_2, X_3 – variable sets of some set algebra, and we will interpret conjunction, disjunction, and negation as intersection, union, and complement in set algebra.

For clarity, let us consider the algebra of all subsets of a set $\{a, b\}$. Then 0 is an empty set, 1 is all sets. Obviously, the sets listed above will be solutions to the original equation in this case, but not only these sets. For example, if we substitute

$X_1 = \{a\}, X_2 = \{a\}, X_3 = \{a\}$, then the equation becomes an identity. From the above, it follows that the method of reduction to perfect disjunctive normal form is not universal for solving Boolean algebra equations in general.

The above considerations show the need for separate consideration of Boolean algebra as a mathematical tool for solving logical equations. Any results obtained for Boolean algebra equations are automatically transferred to FPA equations.

Let us move on to definitions. It is known that a Boolean algebra is any algebra for which the following axioms hold:

$$\begin{aligned} a \vee b &= b \vee a, a \wedge b = b \wedge a, \\ (a \vee b) \vee c &= a \vee (b \vee c), \\ (a \wedge b) \wedge c &= a \wedge (b \wedge c), \\ (a \vee b) \wedge c &= (a \wedge c) \vee (b \wedge c), \\ (a \wedge b) \vee c &= (a \vee c) \wedge (b \vee c), \\ a \vee a &= a, a \wedge a = a, \\ \underline{a} \vee 1 &= 1, a \vee 0 = a, a \wedge 0, a \wedge 1 = a, \\ \bar{a} &= a, \overline{a \vee b} = \bar{a} \wedge \bar{b}, \bar{a} \wedge \bar{b} = \overline{a \vee b} \\ a \vee \bar{a} &= 1, a \wedge \bar{a} = 0. \end{aligned}$$

Let us assume that the algebra is defined on an abstract set, and the operations " $\vee$ ", " $\wedge$ ", " $\bar{}$ " will be called conjunction, disjunction, and negation, respectively.

A Boolean formula is an arbitrary function $f : M^m \rightarrow M$, obtained by superposition of conjunction, disjunction, and negation. Further, for convenience, the conjunction sign will be omitted. The concept of a perfect disjunctive normal form can be defined in exactly the same way as in logic algebra, although it should be noted that in this case the variables in a perfect disjunctive normal form can take values other than 0 and 1.

To effectively solve equations with unknown finite predicates, it is advisable to first study equations written using Boolean algebra operations, since these operations are most commonly encountered when writing predicate equations.

The algebra of finite predicates allows us to interpret knowledge in a rigorous mathematical form, where different features and their meanings are related to each other by Boolean and predicate operations. The classical form of a logical equation with finite predicates looks like this:

$$f(x_1, x_2, \dots, x_n) = 1,$$

where each variable takes values from a finite set of elements, and in general, these domains may be different. In practice, we may encounter problems with many equations.

In this case, the problem can be solved by solving a system of equations:

$$\begin{aligned} f_1(x_1, x_2, \dots, x_n) &= 1, \\ f_2(x_1, x_2, \dots, x_n) &= 1, \\ &\dots \\ f_m(x_1, x_2, \dots, x_n) &= 1. \end{aligned}$$

If necessary, this system can be rewritten as a conjunction of the above equations and presented as a single equation:

$$f_1(x_1, x_2, \dots, x_n) \wedge f_2(x_1, x_2, \dots, x_n) \wedge \dots \wedge f_m(x_1, x_2, \dots, x_n) = 1.$$

Unlike Boolean variables, predicate variables provide greater flexibility in defining the necessary characteristics. For example, consider the following dependencies:

$$\begin{aligned} y^{a_1} &\rightarrow x_1^{b_1} \vee x_1^{b_2}, \\ y^{a_2} &\rightarrow x_1^{b_3} \vee x_1^{b_4}, \\ y^{a_3} &\rightarrow x_1^{b_5} \vee x_1^{b_6}, \end{aligned}$$

where domains for y and x_1 is $\{a_1, a_2, a_3\}$ and $\{b_1, b_2, b_3, b_4, b_5, b_6\}$ accordingly. If $y = a_1$ then $x_1 = b_1$ or $x_1 = b_2$. On the other hand, it follows from the first expression that

$$\neg(x_1^{b_1} \vee x_1^{b_2}) \rightarrow \neg y^{a_1},$$

that means

$$x_1^{b_3} \vee x_1^{b_4} \vee x_1^{b_5} \vee x_1^{b_6} \rightarrow y^{a_2} \vee y^{a_3}.$$

Thus, if x_1 takes on a value from a set $\{b_3, b_4, b_5, b_6\}$ then the property of the object y acquires meaning a_2 or a_3 .

The classic form of a logical equation with finite predicates looks like this:

$$f(x_1, x_2, \dots, x_n) = 1,$$

where each variable takes values from a finite set of elements, and in general, these domains may be different. In practice, you may encounter problems with many equations. In this case, the problem can be solved by solving a system of equations:

$$\begin{aligned} f_1(x_1, x_2, \dots, x_n) &= 1, \\ f_2(x_1, x_2, \dots, x_n) &= 1, \\ &\dots \\ f_m(x_1, x_2, \dots, x_n) &= 1. \end{aligned}$$

If necessary, this system can be rewritten as a conjunction of the above equations and presented as a single equation:

$$f_1(x_1, x_2, \dots, x_n) \wedge f_2(x_1, x_2, \dots, x_n) \wedge \dots \wedge f_m(x_1, x_2, \dots, x_n) = 1.$$

For such equations, we can identify some problems that can be solved:

- 1) find all possible sets of variable values that satisfy this equation. Obviously, this task is difficult because the number of calculations increases exponentially;
- 2) determine whether the system has a solution;
- 3) determine whether it has a unique solution;
- 4) find some important combinations of variable values that satisfy the system;
- 5) solve the system under some initial conditions.

Consider the following system of predicate equations:

$$\begin{aligned} y^{a_1} &\rightarrow g_1(x_1, x_2, \dots, x_n), \\ y^{a_2} &\rightarrow g_2(x_1, x_2, \dots, x_n), \\ &\dots \\ y^{a_m} &\rightarrow g_m(x_1, x_2, \dots, x_n). \end{aligned}$$

This means that when the function y takes on the value a_i , the set of variable values $x_1, x_2, \dots, x_n$ must satisfy the equation

$$g_i(x_1, x_2, \dots, x_n) = 1,$$

which means that if an object has a property a_i , its characteristics $x_1, x_2, \dots, x_n$ must satisfy the above formula.

Generally speaking, the opposite is not necessarily true. If $g_i(x_1, x_2, \dots, x_n) = 1$, y does not necessarily acquire value a_i . Let's consider a stronger dependency:

$$\begin{aligned} y^{a_1} &= g_1(x_1, x_2, \dots, x_n), \\ y^{a_2} &= g_2(x_1, x_2, \dots, x_n), \\ &\dots \\ y^{a_m} &= g_m(x_1, x_2, \dots, x_n). \end{aligned}$$

In this case, any set of feature values $x_1, x_2, \dots, x_n$ either confirms the fact that y is equal to a_i or not (belongs to the corresponding class or not). If an object has a property a_i , then its features must satisfy $g_i(x_1, x_2, \dots, x_n) = 1$. Thus, we can classify an object y by the values of its features $x_1, x_2, \dots, x_n$. It can also be easily shown that the conjunction of any two different functions g_i and g_j is equal to zero. This follows from the basic properties of recognition predicates. It can be shown that the above system is equivalent to the following equation:

$$y^{a_1} g_1(x_1, x_2, \dots, x_n) \vee y^{a_2} g_2(x_1, x_2, \dots, x_n) \vee \dots \vee y^{a_m} g_m(x_1, x_2, \dots, x_n) = 1.$$

Based on the solution of such equations, it is possible to propose logical methods for classifying objects that can be applied to solve practical problems in various areas. Their features can be identified at the stage of constructing a mathematical model, taking into account the characteristics of the data. Usually, propositional logic is used for this purpose. We propose an approach based on finite predicate algebra. Let us construct a general form of such problems based on predicate equations.

6. Means of constructing a system of features for a knowledge base based on a formal algebraic-logical model of a subject area

Logical methods of object classification are used to solve practical problems in various areas: biology, physics, meteorology, etc. Their features can be identified at the stage of constructing a mathematical model, taking into account the characteristics of the data. Usually, propositional logic is used for this purpose. This paper proposes an approach based on finite predicate algebra.

Let us construct a general form of such problems based on predicate equations. Let the feature variables $y_1, y_2, \dots, y_l$ denote certain properties of objects, for example, in the formal description of the symptoms of any disease (according to the diagnosis or treatment protocol) or the assignment of a patient to a certain risk class [15]. Each variable takes values from its domain. Unlike Boolean variables, predicate variables can take values from different arbitrary domains.

Let discrete variables $x_1, x_2, \dots, x_n$ be attributes, based on which we can determine what values variable properties can take. Properties and attributes can be linked together in the form of complex logical dependencies, which can be represented as a predicate equation:

$$P(y_1, y_2, \dots, y_l; x_1, x_2, \dots, x_n) = 1,$$

where P is a finite predicate.

Classifying the object under consideration means determining, based on this predicate equation and experimental data on features $x_1, x_2, \dots, x_n$, which properties (feature values $y_1, y_2, \dots, y_l$) this object possesses and which properties it does not satisfy. For example, each elementary conjunction

$$\begin{aligned} & x_1^{a_{11}} x_2^{a_{12}} \dots x_n^{a_{1n}}, \\ & x_1^{a_{21}} x_2^{a_{22}} \dots x_n^{a_{2n}}, \\ & \dots \\ & x_1^{a_{m1}} x_2^{a_{m2}} \dots x_n^{a_{mn}} \end{aligned}$$

corresponds to its own class of objects. Then, based on a priori dependence $P(y_1, y_2, \dots, y_l; x_1, x_2, \dots, x_n) = 1$ and experimental data on features $x_1, x_2, \dots, x_n$, it is possible to determine which class a given object belongs to. As can be seen from the above considerations, the values of features are grouped into a matrix.

Suppose that as a result of the experiment, we obtained some data related to the values of features $x_1, x_2, \dots, x_n$, describing the classified object and constructed the following predicate equation describing the relationships between them:

$$g(x_1, x_2, \dots, x_n) = 1.$$

The task of classifying objects can be formalized as solving the following predicate equation by finding the unknown predicate f :

$$g(x_1, x_2, \dots, x_n) \rightarrow f(y_1, y_2, \dots, y_l).$$

solving this functional equation, you can determine the values of the features $x_1, x_2, \dots, x_n$, that characterize objects $y_1, y_2, \dots, y_l$.

Unlike Boolean variables, predicate variables provide greater flexibility in defining the necessary characteristics. For example, consider the following dependencies:

$$y^{a_1} \supset x_1^{b_1} \vee x_1^{b_2},$$

$$y^{a_2} \supset x_1^{b_3} \vee x_1^{b_4},$$

$$y^{a_3} \supset x_1^{b_5} \vee x_1^{b_6},$$

where the domains for y and x_1 are $\{a_1, a_2, a_3\}$ and $\{b_1, b_2, b_3, b_4, b_5, b_6\}$ accordingly. If $y = a_1$, then $x_1 = b_1$ or $x_1 = b_2$. On the other hand, it follows from the first expression that

$$\neg(x_1^{b_1} \vee x_1^{b_2}) \supset \neg y^{a_1},$$

that means

$$x_1^{b_3} \vee x_1^{b_4} \vee x_1^{b_5} \vee x_1^{b_6} \supset y^{a_2} \vee y^{a_3}.$$

Thus if x_1 takes a value from the set $\{b_3, b_4, b_5, b_6\}$, the property of object y takes the value a_2 or a_3 .

A universal method for solving FPA systems of equations is to reduce the predicate, given by the system of equations and initial conditions, to a perfect disjunctive normal form (PDNF). However, this procedure involves searching through a large number of intermediate solutions, and its practical implementation requires significant computer time. For some types of linguistic equations, given the

peculiarities of their structure, simpler algorithms for solving them can be developed. For example, a heuristic algorithm is more practical, but it involves searching through all sets of semantic feature values, which slows down the solution process as the number of features increases, with the solution time growing exponentially.

In many practical problems related to knowledge interpretation, which are mostly presented in text form, i.e., related to the semantic processing of natural language information, it is not necessary to obtain all sets of values of semantic features, but it is necessary to obtain one or more sets of feature values (target variables) that are of interest to the user. It is often necessary to find the values of target variables under given initial conditions, which are a fixed set of values of other features. When solving such problems, other variables that are not included in the initial conditions and are not target variables are excluded from the equation by linking them with existential quantifiers.

In research related to logical inferences in knowledge bases, questions arise regarding the determination of the closeness of the relationships between the attributes of these objects, as well as questions about their essentiality and non-essentiality. Obviously, it can be assumed that the formal relationship between features is stronger when fewer sets of values of these variables satisfy the equation. At the same time, if any sets of values of these variables satisfy the original equation, it can be assumed that there is no relationship between these variables.

In addition, the following questions arise when solving practical problems:

1. How will the specific values of this feature, substituted into a logical equation, affect the relationships between other features?
2. How strong is the logical relationship between two (or more) given features?

To answer the first question, it is natural to single out those predicates (and, accordingly, equations) which, when a certain value of the attribute is substituted, turn into predicates that give a stronger relationship between the variables, as well as those predicates, the substitution of this value into which leads to a weakening of the logical relationship between the attributes.

To answer the second question, it is necessary to exclude from the original equation, using the existence quantifier, all variables except those under consideration, and examine the resulting equation with fewer variables, which describes all permissible sets of values of the studied features.

7. Development of an algorithm for eliminating insignificant variables when solving algebraic logic equations to reduce the amount of necessary calculations

Let us consider the feature selection procedure, where the number of features can be reduced. Here we may encounter the following problems.

1. We may need to find some sets of feature values that interest us, where there is at least one value of non-essential features for which there is at least one set of values of essential features. In this case, we apply an existence quantifier to the set of non-essential values.

2. We may need to find some sets of feature values where, for any set of non-essential features, there is at least one solution to the equation. In this case, we apply a universal quantifier to the non-essential variables.

3. We may need to find some sets of attribute values that satisfy the equation, provided that the non-essential attributes take certain specific values.

Let the predicate P depend on variables $x, y, \dots, z$. Let us define the substitution operator $a(P)$ (a belongs to the domain of the variable x), which acts on the predicate P as follows

$$a(P(x, y, \dots, z)) = P(a, y, \dots, z).$$

We will call a substitution operator restrictive if the following condition is met

$$P(a, y, \dots, z) \rightarrow P(x, y, \dots, z)$$

for all $x, y, \dots, z$.

Let's determine the distribution of the substitution operator if the condition is true

$$P(a, y, \dots, z) \leftarrow P(x, y, \dots, z)$$

for all $x, y, \dots, z$.

Interpreting the knowledge presented by this implication, we can say that substitution operators strengthen the logical connection between discrete features, while distributive substitution operators weaken this connection by shifting the relationship between features in an irrelevant manner.

Let us consider the predicate P as follows:

$$P(x, y, \dots, z) = x^{a_1} P_1(y, \dots, z) \vee x^{a_2} P_2(y, \dots, z) \vee \dots \vee x^{a_n} P_n(y, \dots, z).$$

Then

$$a_1(P) = P_1(y, \dots, z) = x^{a_1} P_1(y, \dots, z) \vee x^{a_2} P_1(y, \dots, z) \vee \dots \vee x^{a_n} P_1(y, \dots, z).$$

Obviously, the predicate $a_1(P)$ will be shortened if $P_1 \rightarrow P_i \forall i = 1, 2, \dots, n$.

The operator $a_1(P)$ will distribute if $P_1 \leftarrow P_i \forall i=1,2,\dots,n$.

Let us consider examples of applying the operator a_1 to predicate $P(x,y)$ where variables x,y and z have domains of definition $\{a_1,a_2\}$, $\{b_1,b_2\}$ and $\{c_1,c_2\}$ accordingly.

$$\text{Let } P = x^{a_1}y^{b_1}z^{c_1} \vee x^{a_2}y^{b_1}z^{c_2} \vee x^{a_2}y^{b_1}z^{c_1}.$$

$$\text{Then } a_1(P) = y^{b_1}z^{c_1} = \left(x^{a_1} \vee x^{a_2}\right) \& y^{b_1}z^{c_1} = x^{a_1}y^{b_1}z^{c_1} \vee x^{a_2}y^{b_1}z^{c_1}.$$

With the exception of disjunctions containing the predicate $a_1(P)$, it contains one more disjunction $x^{a_2}y^{b_1}z^{c_1}$, thus the operator a_1 is restrictive for the predicate P . According to the definitions given, in the example above $P_1 = y^{b_1}z^{c_1}$, $P_2 = y^{b_1}z^{c_2} \vee y^{b_1}z^{c_1}$. From this, it is clear that $P_1 \rightarrow P_2$. Let us now consider the predicate

$$\begin{aligned} P &= x^{a_1}y^{b_1}z^{c_1} \vee x^{a_1}y^{b_1}z^{c_2} \vee x^{a_2}y^{b_1}z^{c_1}, \\ a_1(P) &= y^{b_1}z^{c_1} \vee y^{b_1}z^{c_2} = \left(x^{a_1} \vee x^{a_2}\right) \& \left(y^{b_1}z^{c_1} \vee y^{b_1}z^{c_2}\right) = \\ &= x^{a_1}y^{b_1}z^{c_1} \vee x^{a_2}y^{b_1}z^{c_2} \vee x^{a_2}y^{b_1}z^{c_1} \vee x^{a_1}y^{b_1}z^{c_2}. \end{aligned}$$

The operator a_1 for this predicate is obviously distributive. In this example

$$P_1 = y^{b_1}z^{c_1} \vee y^{b_1}z^{c_2} \text{ and } P_2 = y^{b_1}z^{c_2} \text{ thus } P_1 \leftarrow P_2.$$

To answer the second question, it is necessary to exclude all variables except those considered from the original equation and examine the resulting equation with fewer variables, which describes all permissible sets of attribute values. In [51], a fairly broad class of predicates is considered for which an effective algorithm for excluding variables without increasing the size of the original formula can be specified. Here, we extend this class by adding some additional properties. Consider the following properties of the existence quantifier:

1. $\exists x x^a = 1$.
2. $\exists x \neg x^a = 1$.
3. $\exists x \left(\neg (P(x)Q(x)) \right) = \exists x \neg P(x) \vee \exists x \neg Q(x)$.
4. $\exists x (P(x) \vee Q(x)) = \exists x P(x) \vee \exists x Q(x)$.
5. $\exists x (P(x) \& Q(y)) = \exists x P(x) \& Q(y)$.
6. $\exists y (P(x) \rightarrow Q(y)) = P(x) \rightarrow \exists y Q(y)$.

$$7. \exists y(P(x) \rightarrow Q(y)) = P(x) \rightarrow \exists yQ(y).$$

$$8. \text{ Assume } P_i(x) \& P_j(x) = 0, i \neq j, i, j = 1, 2, \dots, k$$

9. Suppose:

$$\begin{aligned} & \exists y((P_1(x) \rightarrow Q_1(y)) \& (P_2(x) \rightarrow Q_2(y)) \& \dots \& (P_k(x) \rightarrow Q_k(y))) = \\ & = (P_1(x) \rightarrow \exists yQ_1(y)) \& (P_2(x) \rightarrow \exists yQ_2(y)) \& \dots \& (P_k(x) \rightarrow \exists yQ_k(y)). \end{aligned}$$

10. If the identity $P_i(x) \equiv 0$ is not true for any $i = 1, 2, \dots, k$ and $P_i(x) \& P_j(x) = 0$ for $i \neq j, i, j = 1, 2, \dots, k$, then

$$\begin{aligned} & \exists x((P_1(x) \rightarrow Q_1(y)) \& (P_2(x) \rightarrow Q_2(y)) \& \dots \& (P_k(x) \rightarrow Q_k(y))) = \\ & = Q_1(y) \vee Q_2(y) \vee \dots \vee Q_k(y) \end{aligned}$$

The above properties allow us to describe a broad class of finite predicates (according to equations) defined on a set of variables $\{x, y, \dots, z\}$, for which it is easy to find relationships between selected variables without increasing the size of the original formulas. Let us define this class recursively.

1. All «recognitions» for variable $x^a, x^b, \dots, x^c$ ($a, b, \dots, c$ – symbols belonging to the domain x) belong to Σ_x .

2. All objections $\neg x^a, \neg x^b, \dots, \neg x^c$ belong to Σ_x .

3. If predicates $\neg P(x), \neg Q(x)$ belong to Σ_x then predicate $\neg(P(x)Q(x))$ belongs to Σ_x .

4. Any predicate that does not depend on a variable x belongs to the set Σ_x .

5. If predicates P_1 та P_2 belong to Σ_x then predicate $P = P_1 \vee P_2$ belongs to Σ_x .

6. If predicate P_1 belongs to Σ_x and predicate P_2 does not depend on x then predicate $P = P_1 \& P_2$ belongs to Σ_x .

7. If predicate P_1 does not depend on x i predicate P_2 belongs to Σ_x then predicate $P = P_1 \rightarrow P_2$ belongs to Σ_x .

8. Let predicates $P_1, P_2, \dots, P_k$ do not depend on x ; $P_i \& P_j = 0$ for $i \neq j, i, j = 1, 2, \dots, k$, predicates $Q_1, Q_2, \dots, Q_k$ belong to Σ_x ; then

$$P = (P_1 \rightarrow Q_1) \& (P_2 \rightarrow Q_2) \& \dots \& (P_k \rightarrow Q_k)$$

belongs to Σ_x .

9. If predicates $P_1, P_2, \dots, P_k$ depend only on x , $P_i \& P_j = 0$ for $i \neq j, i, j = 1, 2, \dots, k$; for any $i = 1, 2, \dots, k$ then identity $P_i \equiv 0$ is not true; predicates $Q_1, Q_2, \dots, Q_k$ do not depend on x ; then predicate

$$P = (P_1 \rightarrow Q_1) \& (P_2 \rightarrow Q_2) \& \dots \& (P_k \rightarrow Q_k)$$

belongs to Σ_x .

It may also be necessary to eliminate redundant variables using a universal quantifier. In this case, we can use the following properties of this quantifier:

1. $\forall x x^a = 0$.
2. $\forall x \neg x^a = 0$.
 $\forall x \neg (P(x) \vee Q(x)) = \forall x \neg P(x) \& \forall x \neg Q(x)$
3. $\forall x (P(x) \& Q(x)) = \forall x P(x) \& \forall x Q(x)$.
4. $\forall x (P(x) \vee Q(y)) = \forall x P(x) \vee Q(y)$.
5. $\forall y (P(x) \& Q(y)) = P(x) \& \forall y Q(y)$.
6. Suppose.

$$P_i(x) \& P_j(x) = 0, i \neq j, i, j = 1, 2, \dots, k,$$

then:

$$\begin{aligned} \forall y ((P_1(x) \& Q_1(y)) \vee (P_2(x) \& Q_2(y)) \vee \dots \vee (P_k(x) \& Q_k(y))) = \\ = (P_1(x) \& \forall y Q_1(y)) \vee (P_2(x) \& \forall y Q_2(y)) \vee \dots \vee (P_k(x) \& \forall y Q_k(y)). \end{aligned}$$

7. If identity $P_i(x) \equiv 0$ is not true for any $i = 1, 2, \dots, k$ and $P_i(x) \& P_j(x) = 0$ for $i \neq j, i, j = 1, 2, \dots, k$ then

$$\begin{aligned} \forall x ((P_1(x) \& Q_1(y)) \vee (P_2(x) \& Q_2(y)) \vee \dots \vee (P_k(x) \& Q_k(y))) = \\ = Q_1(y) \& Q_2(y) \& \dots \& Q_k(y). \end{aligned}$$

We can recursively define a class of predicates Σ_x from which a variable x can be excluded without increasing the size of the formula:

1. All «recognitions» $x^a, x^b, \dots, x^c$ belong to Σ_x .
2. All objections $\neg x^a, \neg x^b, \dots, \neg x^c$ that do not depend on x belong to Σ_x .
3. If $\neg P_1$ and $\neg P_2$ belong to Σ_x then $\neg(P_1 \vee P_2)$ belongs to Σ_x .
4. If predicates P_1 and P_2 belong to Σ_x then predicate $P = P_1 \& P_2$

belongs to Σ_x .

5. If a P_1 belongs to Σ_x and predicate P_2 does not depend on x then predicate $P = P_1 \vee P_2$ belongs to Σ_x .

6. If predicate P_1 does not depend on x and predicate P_2 belongs to Σ_x then predicate $P = P_1 \& P_2$ belongs to Σ_x .

7. Let's assume that predicates $P_1, P_2, \dots, P_k$ do not depend on x , $P_i \& P_j = 0$ for $i \neq j, i, j = 1, 2, \dots, k$; predicates $Q_1, Q_2, \dots, Q_k$ belong to Σ_x , then

$$P = (P_1 \& Q_1) \vee (P_2 \& Q_2) \vee \dots \vee (P_k \& Q_k)$$

belongs to Σ_x .

8. If predicates $P_1, P_2, \dots, P_k$ depend only on x , $P_i \& P_j = 0$ for $i \neq j, i, j = 1, 2, \dots, k$; for any $i = 1, 2, \dots, k$ identity $P_i \equiv 0$ is not true, predicates $Q_1, Q_2, \dots, Q_k$ do not depend on x then predicate $P = (P_1 \& Q_1) \vee (P_2 \& Q_2) \vee \dots \vee (P_k \& Q_k)$ belongs to the set Σ_x .

An example of how this method is used is shown in the next paragraph based on information screening of medical records.

8. Application of algebraic-logical modeling for information screening in conditions of incomplete information

All areas that work with knowledge bases use information technology and produce or consume data. The quality of this data is crucial for the effectiveness of decision-making support processes. Therefore, the relevant areas of IT development are: interpretation of knowledge, its extraction from various sources, intelligent data processing, and the formation of high-quality information for making appropriate decisions. Information quality is an integral characteristic that shows the degree of suitability of data for decision-making. The process of improving information quality, or information screening, is a process of semantic processing of poorly structured, incomplete, and ambiguous data in order to analyze characteristics and identify patterns and inconsistencies [15].

For example, consider the task of analyzing certain medical records, namely data stored in patient medical records, and show that such data is generally characterized by the following features: incompleteness, obsolescence, inconsistency, etc. This situation leads to more than one interpretation, i.e., to ambiguity of information. The use of the results of processing ambiguous information contained in medical documentation with low data quality leads to incorrect management and medical decisions. In this regard, it is important to introduce models and methods

for processing medical data that improve the quality of information for decision-making, namely: improve data completeness by extracting additional information, improve data accuracy by using only relevant information, remove contradictory information, and form a relevant data set.

Research methods for this type of task are based on the use of system analysis methods, decision-making theory, intelligent data analysis methods, and intelligence theory methods. Namely: the algebraic-logical modeling method and the comparator identification method are used in the intelligent processing of data from patient medical records for information screening of medical documentation; the mathematical apparatus of optimization methods is used in solving the problem of planning therapeutic and preventive measures based on the technology of information screening of medical documentation; a service-oriented approach was used to develop tools for solving the tasks set.

In works [15, 53], a method for information screening of medical documentation was proposed and developed. The use of the comparator identification method for solving the problem of information screening of medical documentation was also proposed. A model for identifying medical and diagnostic parameters has been developed.

For example, to improve the quality of information for medical decision-making, a method of information screening has been formulated and justified, which consists of using a set of data processing models (Fig. 1).

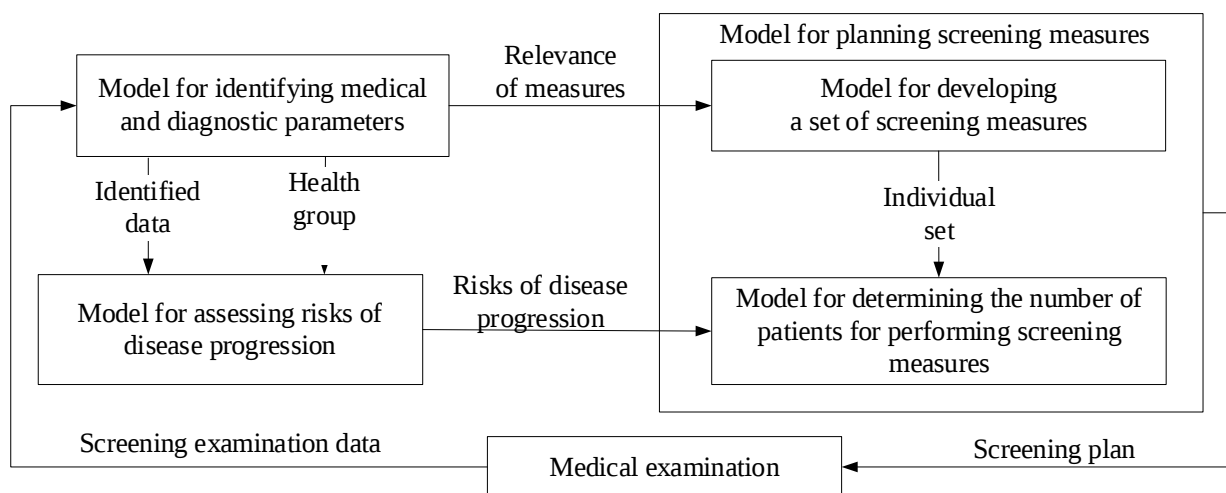


Fig. 1. Medical documentation information screening method

To solve the problem of information screening of certain medical documentation (for example, medical records filled out by a family doctor), the use of an algebraic-logical method of modeling a system of features based on finite

predicate algebra is proposed as a method for modeling subject knowledge. The general scheme for using the proposed method is as follows: based on information from medical records and other sources of medical information, a set of features is formed, which is modeled by predicate equations. Depending on the results to be obtained, the reactions of the predicates are subjected to semantic processing and grouped into a set of aggregated signs, which subsequently form the basis for medical decisions (Fig. 1). In this formulation of the problem, the comparator is a system of predicate equations, the solution of which is described above.

One of the main conditions for the application of the comparator identification method is the discreteness, finiteness, and determinism of the objects of the subject space. In this case, the subject space U is a Cartesian product of a set of objects from the subject space $U_1 \times U_2 \times \dots \times U_n$: medical data, patient records, clinical monitoring data, etc. With the help of a predicate $P(x_1, x_2, \dots, x_m)$ any ratio that is defined at the subject space $U = U_1 \times U_2 \times \dots \times U_n$ is described. In the language of finite predicate algebra, this predicate is formalized using the recognition predicate and the basic operations of conjunction and disjunction. The recognition predicate models the ability of a person to unambiguously assign the object under consideration to one of two classes. The predicate is equal to 1 when the object is recognized and 0 otherwise. The introduced predicate variables are linked by logical equations, the joint solution of which assigns the elements of the set of objects under consideration to a certain class, i.e., determines the patient's risk group for the presence of certain diseases.

The following experiment plan is proposed. We use real medical data and encode it using predicate equations. Note that although some variables may take on the value «unknown», this is still a closed-world case, since «unknown» means only a value from the alphabet in which the variable is defined. Each domain for any variable is closed. After we have written a system of equations with the help of experts, we begin to remove variables that we consider irrelevant at the moment. This does not mean that in other cases other variables will be considered irrelevant. Important variables are those for which we want to determine logical relationships. The result is an equation from which irrelevant variables have been removed. The resulting equation is simpler than the original system, and it is easier to analyze the relationships between relevant variables.

If we consider information screening of medical data for the assessment of the development and prevention of cardiovascular diseases [22], we can identify a set of features for formalizing screening procedures. Let us consider the following features and their meanings:

Gender: $X_1 = \{x_1^1, x_1^2\}$, where x_1^1 means female, x_1^2 means male.

Age: $X_2 = \{x_2^1, x_2^2, x_2^3\}$, where x_2^1 under 40 years, x_2^2 from 40 to 50 years, x_2^3 over 50 years.

Diabetes mellitus: $X_3 = \{x_3^1, x_3^2, x_3^3, x_3^4\}$, where x_3^1 – yes, x_3^2 – no (actual diagnosis), x_3^3 – no (diagnosis not determined), x_3^4 – unknown.

Arterial hypertension: $X_4 = \{x_4^1, x_4^2, x_4^3, x_4^4\}$, where x_4^1 – yes, x_4^2 – no (actual diagnosis), x_4^3 – no (diagnosis not determined), x_4^4 – unknown.

Kidney problems: $X_5 = \{x_5^1, x_5^2, x_5^3\}$ where x_5^1 – yes, x_5^2 – no, x_5^3 – unknown.

Tachycardia: $X_6 = \{x_6^1, x_6^2, x_6^3, x_6^4, x_6^5\}$, where x_6^1 – yes (actual diagnosis), x_6^2 – yes (diagnosis not determined), x_6^3 – no (actual diagnosis), x_6^4 – no (diagnosis not determined), x_6^5 – unknown.

Heredity in relation to heart and vascular diseases: $X_7 = \{x_7^1, x_7^2, x_7^3\}$, where x_7^1 – yes, x_7^2 – no, x_7^3 – unknown.

Smoking: $X_8 = \{x_8^1, x_8^2, x_8^3\}$, where x_8^1 – yes, x_8^2 – no, x_8^3 – unknown.

Alcohol problems: $X_9 = \{x_9^1, x_9^2, x_9^3\}$, where x_9^1 – yes, x_9^2 – no, x_9^3 – unknown.

Hypodynamia: $X_{10} = \{x_{10}^1, x_{10}^2, x_{10}^3\}$, where x_{10}^1 – yes, x_{10}^2 – no, x_{10}^3 – unknown.

These features make it possible to develop a model for identifying diagnostic parameters, which can be used to determine the patient's health group $R = \{r_1, r_2, r_3, r_4\}$ where r_1 low risk of heart and vascular disease, r_2 medium risk, r_3 high risk, r_4 very high risk.

A set of aggregated characteristics $Q_1 - Q_3$ is used to determine the health group, where Q_1 is denoted by X_1 and X_2 , Q_2 is denoted by X_7 to X_{10} , Q_3 is denoted by $X_3 - X_6$.

The values for each health group and each aggregated indicator are divided into four classes in accordance with the relevant medical and technological documentation (standardized clinical protocol and local protocols relating to the prevention of heart and vascular diseases).

The system of equations for forming, for example, a feature Q_2 looks like this:

$$\left\{ \begin{aligned} q_2^1 &= x_7^2 x_8^2 \left(x_9^2 \vee x_9^3 \left(x_{10}^2 \vee x_{10}^3 \right) \right) \vee x_7^2 x_8^3 x_9^2 x_{10}^2 \vee x_7^3 x_8^2 x_{10}^2 \left(x_9^2 \vee x_9^3 \right), \\ q_2^2 &= x_7^2 \left(x_8^1 \left(x_9^1 x_{10}^2 \vee x_9^2 \right) \vee x_9^3 \left(x_8^1 x_{10}^2 \vee x_8^2 x_{10}^1 \right) \right) \vee \\ &\vee \left(x_7^2 \left(x_8^2 x_9^1 \vee x_8^3 x_9^2 \right) \vee \left(x_7^2 x_9^3 \vee x_7^3 x_9^1 \right) x_8^3 x_{10}^2 \vee \right. \\ &\vee \left(x_7^2 x_8^3 \vee x_7^3 x_8^2 \right) x_9^1 \left(x_{10}^2 \vee x_{10}^3 \right) \vee x_7^3 x_8^2 \left(x_9^2 \vee x_9^3 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee \\ &\vee x_7^3 x_8^1 \left(x_9^2 x_{10}^2 \vee x_9^3 \right) \vee x_7^3 x_8^3 x_9^2 \left(x_{10}^1 \vee x_{10}^2 \right), \\ q_2^3 &= x_7^1 x_{10}^2 \left(x_8^1 x_9^2 \vee x_8^2 \left(x_9^1 \vee x_9^2 \right) \right) \vee \\ &\vee \left(x_7^1 x_9^3 \left(x_8^1 \vee x_8^2 \right) \vee \left(x_7^1 x_8^3 \vee x_7^3 x_8^1 \right) x_9^1 \right) \left(x_{10}^2 \vee x_{10}^3 \right) \vee x_7^1 x_8^3 \left(x_9^2 \vee x_9^3 \right) \vee \\ &\vee \left(x_7^2 \left(x_8^1 \left(x_9^1 \vee x_9^3 \right) \vee x_8^3 x_9^3 \right) \vee \left(x_7^2 x_8^3 \vee x_7^3 x_8^2 \right) x_9^1 x_{10}^1 \vee \right. \\ &\vee x_7^3 \left(x_8^1 x_9^2 \vee x_8^3 x_9^1 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_7^3 x_8^3 \left(x_9^2 x_{10}^3 \vee x_9^3 \right), \\ q_2^4 &= x_7^1 x_9^3 x_{10}^1 \left(x_8^1 \vee x_8^2 \right) \vee \left(x_7^1 x_8^3 \vee x_7^3 x_8^1 \right) x_9^1 x_{10}^1 \vee \\ &\vee \left(x_7^1 x_8^2 x_9^1 \vee x_7^1 x_9^2 \left(x_8^1 \vee x_8^2 \right) \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_7^1 x_8^1 x_9^1. \end{aligned} \right.$$

A model for identifying medical and diagnostic parameters using the example of risk class determination is presented in the form of a system of predicate equations:

$$\left\{ \begin{array}{l} r_1 = q_1^1 q_2^1 (q_3^1 \vee q_3^2) \vee (q_1^1 q_2^2 \vee (q_1^2 \vee q_1^3) q_2^1) q_3^1, \\ r_2 = q_1^1 (q_2^1 q_3^3 \vee q_2^2 q_3^2) \vee \\ \quad \vee (q_1^1 (q_2^3 \vee q_2^4) \vee q_1^2 (q_2^2 \vee q_2^3) \vee q_1^3 q_2^2 \vee q_1^4 (q_2^1 \vee q_2^2)) (q_3^1 \vee q_3^2) \vee \\ \quad \vee (q_1^2 \vee q_1^3) q_2^1 (q_3^2 \vee q_3^3) \vee (q_1^2 q_2^4 \vee (q_1^3 \vee q_1^4) q_2^3) q_3^1, \\ r_3 = q_2^1 q_3^4 \vee (q_1^1 \vee q_1^2 \vee q_1^3) (q_2^2 \vee q_2^3) (q_3^3 \vee q_3^4) \vee \\ \quad \vee q_1^3 q_3^2 ((q_2^3 \vee q_2^4) \vee (q_1^3 \vee q_1^4) q_2^4 q_3^1 \vee \\ \quad \vee (q_1^1 q_2^4 \vee q_1^4 (q_2^1 \vee q_2^2)) q_3^3 \vee (q_1^2 q_2^4 \vee q_1^4 q_2^3) (q_3^2 \vee q_3^3), \\ r_4 = (q_1^1 \vee q_1^2) q_2^4 q_3^4 \vee q_1^3 q_2^4 (q_3^3 \vee q_3^4) \vee \\ \quad \vee q_1^4 q_3^4 (q_2^2 \vee q_2^3) \vee q_1^4 q_2^4 (q_3^2 \vee q_3^3 \vee q_3^4). \end{array} \right.$$

The final classification can be described by the following system:

$$\left\{ \begin{array}{l} r_1 = q_1^1 q_2^1 (q_3^1 \vee q_3^2) \vee (q_1^1 q_2^2 \vee (q_1^2 \vee q_1^3) q_2^1) q_3^1, \\ r_2 = q_1^1 (q_2^1 q_3^3 \vee q_2^2 q_3^2) \vee \\ \quad \vee (q_1^1 (q_2^3 \vee q_2^4) \vee q_1^2 (q_2^2 \vee q_2^3) \vee q_1^3 q_2^2 \vee q_1^4 (q_2^1 \vee q_2^2)) (q_3^1 \vee q_3^2) \vee \\ \quad \vee (q_1^2 \vee q_1^3) q_2^1 (q_3^2 \vee q_3^3) \vee (q_1^2 q_2^4 \vee (q_1^3 \vee q_1^4) q_2^3) q_3^1, \\ r_3 = q_2^1 q_3^4 \vee (q_1^1 \vee q_1^2 \vee q_1^3) (q_2^2 \vee q_2^3) (q_3^3 \vee q_3^4) \vee \\ \quad \vee q_1^3 q_3^2 (q_2^3 \vee q_2^4) \vee (q_1^3 \vee q_1^4) q_2^4 q_3^1 \vee \\ \quad \vee (q_1^1 q_2^4 \vee q_1^4 (q_2^1 \vee q_2^2)) q_3^3 \vee (q_1^2 q_2^4 \vee q_1^4 q_2^3) (q_3^2 \vee q_3^3), \\ r_4 = (q_1^1 \vee q_1^2) q_2^4 q_3^4 \vee q_1^3 q_2^4 (q_3^3 \vee q_3^4) \vee q_1^4 q_3^4 (q_2^2 \vee q_2^3) \vee q_1^4 q_2^4 (q_3^2 \vee q_3^3 \vee q_3^4) \end{array} \right.$$

Let's explore the logical connections between discrete features $x_1 - x_{10}$. First, let's rewrite the system of predicate equations in the following form:

$$\begin{aligned}
P(q_2, x_1, \dots, x_{10}) = & \\
= q_2^1 & \left(x_7^2 x_8^2 \left(x_9^2 \vee x_9^3 \left(x_{10}^2 \vee x_{10}^3 \right) \right) \vee x_7^2 x_8^3 x_9^2 x_{10}^2 \vee x_7^3 x_8^2 x_{10}^2 \left(x_9^2 \vee x_9^3 \right) \right) \vee \\
\vee q_2^2 & \left(x_7^2 \left(x_8^1 \left(x_9^1 x_{10}^2 \vee x_9^2 \right) \vee x_9^3 \left(x_8^1 x_{10}^2 \vee x_8^2 x_{10}^1 \right) \right) \vee \right. \\
\vee & \left(x_7^2 \left(x_8^2 x_9^1 \vee x_8^3 x_9^2 \right) \vee \left(x_7^2 x_9^3 \vee x_7^3 x_9^1 \right) x_8^3 x_{10}^2 \vee \right. \\
\vee & \left(x_7^2 x_8^3 \vee x_7^3 x_8^2 \right) x_9^1 \left(x_{10}^2 \vee x_{10}^3 \right) \vee x_7^3 x_8^2 \left(x_9^2 \vee x_9^3 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee \\
\vee & x_7^3 x_8^1 \left(x_9^2 x_{10}^2 \vee x_9^3 \right) \vee x_7^3 x_8^3 x_9^2 \left(x_{10}^1 \vee x_{10}^2 \right) \left. \right) \vee \\
\vee q_2^3 & \left(x_7^1 x_{10}^2 \left(x_8^1 x_9^2 \vee x_8^2 \left(x_9^1 \vee x_9^2 \right) \right) \vee \right. \\
\vee & \left(x_7^1 x_9^3 \left(x_8^1 \vee x_8^2 \right) \vee \left(x_7^1 x_8^3 \vee x_7^3 x_8^1 \right) x_9^1 \right) \left(x_{10}^2 \vee x_{10}^3 \right) \vee x_7^1 x_8^3 \left(x_9^2 \vee x_9^3 \right) \vee \\
\vee & \left(x_7^2 \left(x_8^1 \left(x_9^1 \vee x_9^3 \right) \vee x_8^3 x_9^3 \right) \vee \left(x_7^2 x_8^3 \vee x_7^3 x_8^2 \right) x_9^1 x_{10}^1 \vee \right. \\
\vee & x_7^3 \left(x_8^1 x_9^2 \vee x_8^3 x_9^1 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_7^3 x_8^3 \left(x_9^2 x_{10}^3 \vee x_9^3 \right) \left. \right) \vee \\
\vee q_2^4 & \left(x_7^1 x_9^3 x_{10}^1 \left(x_8^1 \vee x_8^2 \right) \vee \left(x_7^1 x_8^3 \vee x_7^3 x_8^1 \right) x_9^1 x_{10}^1 \vee \right. \\
\vee & \left(x_7^1 x_8^2 x_9^1 \vee x_7^1 x_9^2 \left(x_8^1 \vee x_8^2 \right) \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_7^1 x_8^1 x_9^1 \left. \right) = 1
\end{aligned}$$

It is clear that this predicate belongs to the class Δ_{x7} . Let us examine the relationship between all variables except x_7 . This exclusion will give us the relationship between the variables $q_2, x_1, \dots, x_6, x_8, x_9, x_{10}$:

$$\begin{aligned}
F &= \exists x_7 P(q_2, x_1, \dots, x_{10}) = \\
&= q_2^1 \left(x_8^2 \left(x_9^2 \vee x_9^3 \left(x_{10}^2 \vee x_{10}^3 \right) \right) \vee x_8^3 x_9^2 x_{10}^2 \vee x_8^2 x_{10}^2 \left(x_9^2 \vee x_9^3 \right) \right) \vee \\
&\vee q_2^2 \left(x_8^1 \left(x_9^1 x_{10}^2 \vee x_9^2 \right) \vee x_9^3 \left(x_8^1 x_{10}^2 \vee x_8^2 x_{10}^1 \right) \right) \vee \left(\left(x_8^2 x_9^1 \vee x_8^3 x_9^2 \right) \vee \left(x_9^3 \vee x_9^1 \right) x_8^3 x_{10}^2 \vee \right. \\
&\vee \left(x_8^3 \vee x_8^2 \right) x_9^1 \left(x_{10}^2 \vee x_{10}^3 \right) \vee x_8^2 \left(x_9^2 \vee x_9^3 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_8^1 \left(x_9^2 x_{10}^2 \vee x_9^3 \right) \vee \\
&\vee x_8^3 x_9^2 \left(x_{10}^1 \vee x_{10}^2 \right) \vee \\
&\vee q_2^3 \left(x_{10}^2 \left(x_8^1 x_9^2 \vee x_8^2 \left(x_9^1 \vee x_9^2 \right) \right) \vee \left(x_9^3 \left(x_8^1 \vee x_8^2 \right) \vee \left(x_8^3 \vee x_7^3 x_8^1 \right) x_9^1 \right) \left(x_{10}^2 \vee x_{10}^3 \right) \vee \right. \\
&\vee x_8^3 \left(x_9^2 \vee x_9^3 \right) \vee \left(\left(x_8^1 \left(x_9^1 \vee x_9^3 \right) \vee x_8^3 x_9^3 \right) \vee \left(x_8^3 \vee x_8^2 \right) x_9^1 x_{10}^1 \vee \right. \\
&\vee \left(x_8^1 x_9^2 \vee x_8^3 x_9^1 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_8^3 \left(x_9^2 x_{10}^3 \vee x_9^3 \right) \left. \right) \vee \\
&\vee q_2^4 \left(x_9^3 x_{10}^1 \left(x_8^1 \vee x_8^2 \right) \vee \left(x_8^3 \vee x_8^1 \right) x_9^1 x_{10}^1 \vee \right. \\
&\vee \left(x_8^2 x_9^1 \vee x_9^2 \left(x_8^1 \vee x_8^2 \right) \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_8^1 x_9^1 \left. \right) = 1
\end{aligned}$$

It should be noted that the size of the output formula has not increased, which is explained by the fact that the predicate $P(q_2, x_1, \dots, x_{10})$ belongs to Σ_{x7} .

Suppose we are interested in the link between q_2, x_9, x_{10} . We remove other features $F(x_1, \dots, x_{10})$ from the predicate:

$$\begin{aligned}
G(q_2, x_9, x_{10}) &= \exists x_1 \exists x_2 \exists x_3 \exists x_4 \exists x_5 \exists x_6 \exists x_7 \exists x_8 P(q_2, x_1, \dots, x_{10}) = \\
&= q_2^1 \left(\left(x_9^2 \vee x_9^3 \left(x_{10}^2 \vee x_{10}^3 \right) \right) \vee x_9^2 x_{10}^2 \vee x_{10}^2 \left(x_9^2 \vee x_9^3 \right) \right) \vee \\
&\vee q_2^2 \left(\left(x_9^1 x_{10}^2 \vee x_9^2 \right) \vee x_9^3 \left(x_{10}^2 \vee x_{10}^1 \right) \right) \vee \left(\left(x_9^1 \vee x_9^2 \right) \vee \left(x_9^3 \vee x_9^1 \right) x_{10}^2 \vee \right. \\
&\vee x_9^1 \left(x_{10}^2 \vee x_{10}^3 \right) \vee \left(x_9^2 \vee x_9^3 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee \left(x_9^2 x_{10}^2 \vee x_9^3 \right) \vee \\
&\vee x_9^2 \left(x_{10}^1 \vee x_{10}^2 \right) \vee q_2^3 \left(x_{10}^2 \left(x_9^1 \vee x_9^2 \right) \vee x_9^3 \left(x_{10}^2 \vee x_{10}^3 \right) \vee \left(x_9^2 \vee x_9^3 \right) \vee \right. \\
&\vee \left(\left(x_9^1 \vee x_9^3 \right) \vee x_9^1 x_{10}^1 \vee \left(x_9^2 \vee x_9^1 \right) \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee \\
&\vee \left(x_9^2 x_{10}^3 \vee x_9^3 \right) \vee q_2^4 \left(x_9^3 x_{10}^1 \vee x_9^1 x_{10}^1 \vee \left(x_9^1 \vee x_9^2 \right) \left(x_{10}^1 \vee x_{10}^3 \right) \vee x_9^1 \right) = 1
\end{aligned}$$

We reduced the initial formula and obtained a simpler dependence between the selected characteristics. Once the necessary dependence has been obtained, we can solve the resulting equation with one or more target variables.

Depending on the solution of the system of predicate equations, the medical record will be assigned to a class, and the doctor will develop a set of therapeutic and preventive procedures and recommendations to maintain the patient's health in good condition.

Although some variables in the medical example can take on the value «unknown», we are dealing with a closed world, since «unknown» only means an element from the domain for the variable. All domains are strictly defined and cannot be supplemented by any other elements. The main advantage of this method, based on the specific structure of predicate systems, is that after removing non-essential variables, the original system (or equation) is simplified, which is related to the special properties of quantifiers. Variables that are essential are not necessarily fixed forever. The researcher decides which logical relationships are important at a given moment.

As a further direction of research, it is possible to extend the classes of predicates from which non-essential variables can be easily removed. Such classes are much more complex than relational structures, but many practical tasks require such a representation and analysis of knowledge.

Conclusions

1. It has been shown that finite predicate algebra (FPA) is a universal mathematical apparatus for the formal description of deterministic, discrete, finite objects of various natures. FPA can interpret knowledge in a strict mathematical form, where different features and their meanings are linked by Boolean and predicate operations.

2. It has been shown that when constructing a feature system for a knowledge base, it is advisable to use an algebraic-logical model of the subject area. An approach based on the use of the comparator identification method and FPA is proposed.

3. It is shown that in practical tasks related to the interpretation of knowledge, which is mostly presented in textual form, it is not necessary to obtain all sets of values of semantic features, but it is necessary to obtain one or several essential sets of feature values (target variables), which are a fixed set of values of other features. When solving such problems, other variables that are not included in the initial conditions and are not target variables are excluded from the equation by binding them with existential quantifiers.

4. An algorithm for excluding non-essential variables when solving algebraic-logical equations is proposed, which reduces the number of necessary calculations. At the same time, it is proposed to single out those predicates (and, accordingly, equations) which, when a certain value of a feature is substituted, turn into predicates that give a stronger connection between variables, as well as those predicates, the substitution of this value into which leads to a weakening of the logical connection between features.

5. Finite predicate equations of various types are considered. A description of classification problems based on predicate equations is given. The problem of classifying objects based on features that take discrete values is described mathematically as a solution of predicate equations.

6. A broad class of predicates is described from which redundant variables can be removed and focus can be placed on the relationships between essential variables. A method for removing non-essential variables by applying an existential quantifier is proposed and demonstrated on a real medical example. The example demonstrates the use of the mathematical apparatus of intelligence theory, the method of comparator identification, and the tools of finite predicate algebra on a model of identification of medical diagnostic parameters in the form of a system of predicate equations, the solution of which provides an interpretation of medical knowledge in a specific field.

References

1. Shabanov-Kushnarenko, Yu.P. Theory of Intelligence. Mathematical Means. Kharkiv: Higher School. Published by Kharkiv University, 1984. 144 pp.
2. Shabanov-Kushnarenko, Yu.P. Theory of Intelligence. Technical Means. Kharkiv: Higher School. Published by Kharkiv University, 1986. – 136 p.
3. Shabanov-Kushnarenko Yu.P. Theory of Intelligence. Problems and Prospects. –Kharkiv: Higher School. Published by Kharkiv University, 1987. – 159 p.
4. Bondarenko M.F., Shabanov-Kushnarenko Yu.P. Theory of Intelligence: Textbook. – Kharkiv: LLC «SMIT Company», – 2006. – 576 p.
5. Bondarenko M.F., Shabanov-Kushnarenko Yu.P. Brain-like Structures. Reference Book. Volume One. Edited by Academician I.V. Sergienko, National Academy of Sciences of Ukraine. – Kiev, Naukova Dumka, 2011. – 460 pp.
6. Bondarenko M.F., Shabanov-Kushnarenko Yu.P., Sharonova N.V. Tools for Comparative Identification. – Bionics of Intelligence. – 2010, No. 2(73), pp. 74–86.
7. Bondarenko M.F., Shabanov-Kushnarenko Yu.P., Sharonova N.V. Situational-textual predicate. – Bionics of Intelligence. – 2010, No. 2(73), pp. 87–98.
8. Bondarenko M.F., Shabanov-Kushnarenko Yu.P., Sharonova N.V. Boolean structure of text. – Bionics of Intelligence. – 2010, No. 2(73), pp. 99–110.
9. Bulkin V.I., Sharonova N.V. Mathematical models of knowledge and their implementation using algebraic predicate structures: monograph // Bulkin V.I., Sharonova N.V. – NTU «KhPI», MEGI. – Donetsk, 2010. – 304 p.
10. Kanishcheva O.V., Sharonova N.V. Linguistic technologies for knowledge identification in information systems. Linguistic technologies for knowledge extraction and processing by intelligent systems. – LAP LAMBERT Academic Publishing? – 2013. – 164 p.
11. Khayrova N.F., Sharonova N.V. Information and linguistic technologies for extraction and identification of deep knowledge in texts: monograph // N.F. Khayrova, N.V. Sharonova. – Kharkiv: FLP Koryak S.F., 2016. – 205 p.
12. Karatayev O.A., Shubin I.Yu. Problems of knowledge reuse in software system design / O. A. Karatayev, I.Yu. Shubin // Current state of scientific research and technologies in industry, 2023. No. 2(24). – pp. 25–34. DOI: <https://doi.org/10.30837/ITSSI.2023.24.025>
13. Karataev O.A., Sitnikov D.E. Method for reusing knowledge in the form of logical equations / O.A. Karataev, D.E. Sitnikov // Current state of scientific research and technologies in industry, 2023. No. 3(25). – pp. 25–34. DOI: <https://doi.org/10.30837/ITSSI.2023.24.025>
14. Karataiev O., Sitnikov D., Sharonova N. A Method for Investigating Links between Discrete Data Features in Knowledge Bases in the Form of Predicate Equations. CEUR Workshop Proceedings, Vol. 3387. 2023. pp. 224–235. URL: <https://ceur-ws.org/Vol-3387/paper17.pdf>
15. Melnik K. Towards medical screening information technology: the healthgrid-based approach/ K. Melnik, O. Cherednichenko, V. Glushko // Information Systems: Methods, Models, and Applications. – Heidelberg : Springer, 2013. – P. 202–204.
16. Olga Cherednichenko, Nataliia Sharonova, Ganna Pliekhova, Nadiia Babkova: Intelligent Methods of Secure Routing in Software-Defined Networks. Proceedings of the 8th International Conference on Computational Linguistics and Intelligent Systems. Volume I: Machine Learning Workshop, Lviv, Ukraine, April 12–13, 2024. CEUR Workshop Proceedings 3664, CEUR-WS.org 2024 Pp.342–351.

17. Victoria Vysotska, Kirill Smelyakov, Natalia Sharonova, Maksym Derenskyi, Ganna Pliekhova and Vadim Repikhov. Information System for Monitoring and Planning Maintenance of Offshore Wind Farms - Conference on Computational Linguistics and Intelligent Systems. Volume II: Modeling, Optimization, and Controlling in Information and Technology Systems Workshop 2024, Lviv, Ukraine, April 12-13, 2024. CEUR Workshop Proceedings 3664, CEUR-WS.org 2024 Pp. 63–82.
18. Cui, X., Tang, D., Zhu, H., & Yin, L. (2017). Knowledge representation and semantic inference of process based on ontology and swrl. *Machine Building & Automation*, 46(3), 6–10
19. Moshood, T. D., Rotimi, F. E., Rotimi, J. O. B. 2022. An Integrated Paradigm for Managing Efficient Knowledge Transfer: Towards a More Comprehensive Philosophy of Transferring Knowledge in the Construction Industry. *Construction Economics and Building*, 22:3, 65–98. DOI: <https://doi.org/10.5130/AJCEB.v22i3.8050>
20. Z. Chen, Y. Wang, B. Zhao, J. Cheng, X. Zhao, Z. Duan, Knowledge graph completion: A review, *IEEE Access* 8 (2020) 192435–192456
21. Schacht, S., Maedche, A. (2016). A Methodology for Systematic Project Knowledge Reuse. *Innovations in Knowledge Management*. Razmerita, L. (Eds.), Springer (Berlin), 2016, pp. 19–44.
22. Ameri, F., Dutta, D. (2005). Product lifecycle management: closing the knowledge loops. *Computer-Aided Design and Applications* 2 (5), pp. 577–590.
23. Milton, N. R. (2007). *Knowledge acquisition in practice: a step-by-step guide*. Springer Science & Business Media.
24. Carro Saavedra C., Serrano Villodres T., Lindemann U. (2016). Review and Classification of Knowledge in Engineering Design. *IcoRD2017*. Technische Universität München
25. Philippe Martin and Jérémy Bénard. 2016. Top-level Ideas about Importing, Translating and Exporting Knowledge via an Ontology of Representation Languages. In *Proceedings of the 12th International Conference on Semantic Systems (SEMANTiCS 2016)*. Association for Computing Machinery, New York, NY, USA, 89–92. DOI: <https://doi.org/10.1145/2993318.2993344>
26. Khudhair, A. T. (2017). The intelligence theory mathematical apparatus formal base. *Advanced Information Systems*, 1(1), 38–43. DOI: <https://doi.org/10.20998/2522-9052.2017.1.07>
27. Sharonova, Natalia et al. «Issues of Fact-based Information Analysis». *International Conference on Computational Linguistics and Intelligent Systems* (2018).
28. P. G. Omran, K. Wang, Z. Wang, An Embedding-based Approach to Rule Learning in Knowledge Graphs, *TKDE* 33 (2021) 1348–1359.
29. 13 T. Pellissier-Tanon, G. Weikum, F. Suchanek, YAGO 4: A Reasonable Knowledge Base, in: *ESWC*, 2020
30. Omran, P. G., Wang, Z., & Wang, K. (2022). Learning Rules With Attributes and Relations in Knowledge Graphs. In *AAAI Spring Symposium: MAKE*.
31. P. G. Omran, K. Wang, Z. Wang, Scalable rule learning via learning representation, in: *IJCAI*, 2018.
32. M.Svato, S.Schockaert, J.Davis, STRiKE: Rule-Driven Relational Learning Using Stratified k-Entailment, in: *ECAI*, 2020.
33. G. A. Oparin, V. G. Bogdanova, and A. A. Pashinin, «Classification in a binary feature space using logical dynamic models», 44th International Congress on Information, Communication and Electronic Technologies (MIPRO), Opatija, Croatia, 2021, pp. 1020–1025, DOI: [10.23919/MIPRO52101.2021.9596697](https://doi.org/10.23919/MIPRO52101.2021.9596697)

34. A. Kabulov, E. Urunboev, and I. Saimanov, «Method for object recognition based on logical correction functions,» 2020 International Conference on Information Sciences and Communication Technologies (ICISCT), Tashkent, Uzbekistan, 2020, pp. 1–4, doi: 10.1109/ICISCT50599.2020.9351473. pp. 1–4, DOI: 10.1109/ICISCT50599.2020.9351473
35. M. Ahmed, S. K. Ibrahim, and S. Yakut, «Classification of hyperspectral images based on logical data analysis,» 2019 IEEE Aerospace Conference, Big Sky, MT, USA, 2019, pp. 1–9, DOI: 10.1109/AERO.2019.8742023
36. S. Kim, S. Noh and H. S. Ryoo, «Identifying Combinatorial Significance for Classification of Alzheimer’s Disease Proteomics Expression with Logical Analysis of Data», 2021 IEEE International Conference on Bioinformatics and Biomedicine (BIBM), Houston, Texas, USA, 2021, pp. 1661–1663, DOI: 10.1109/BIBM52615.2021.9669835
37. Povkhan and M. Lupe, «Algorithmic Classification Trees,» 2020 IEEE Third International Conference on Intelligent Data Analysis and Processing (DSMP), Lviv, Ukraine, 2020, pp. 37–43, DOI: 10.1109/DSMP47368.2020.9204198
38. M. Tamilselvi, G. Ramkumar, G. Anita, P. Nirmal, and S. Ramesh, «A New Text Recognition Scheme Using a Classification-Based Digital Image Processing Strategy,» 2022 International Conference on Advances in Computing, Communication, and Applied Informatics (ACCAI), Chennai, India, 2022, pp. 1–6, DOI: 10.1109/ACCAI53970.2022.9752542
39. S. N. Veigl and B. Kovalerchuk, «An interactive visual approach to self-service data classification for democratizing machine learning,» 2020 24th International Conference on Information Visualization (IV), Melbourne, Australia, 2020, c. 280–285, DOI: 10.1109/IV51561.2020.00052
40. B. Jean-Marc and O. Lafitte, «Combining weak classifiers: a logical analysis,» 2021 23rd International Symposium on Symbolic and Numerical Algorithms for Scientific Computing (SYNASC), Timisoara, Romania, 2021, pp. 178–181, DOI: 10.1109/SYNASC54541.2021.00038
41. Z. Chen, A. Xu, Y. Zhou, and Y. Gai, «Research on Pulsar Classification Based on Machine Learning,» 2020 3rd International Conference on Intelligent Autonomous Systems (ICoIAS), Singapore, 2020, pp. 14–18, DOI: 10.1109/ICoIAS49312.2020.9081836
42. Nwankwo, H. Wimmer, L. Chen, and J. Kim, «Text Classification of Digital Forensic Data,» 2020 11th Annual IEEE Conference on Information Technology, Electronics, and Mobile Communications (IEMCON), Vancouver, British Columbia, Canada, 2020, pp. 0661–0667, DOI: 10.1109/IEMCON51383.2020.9284913
43. Shubin, S. Snisar, and S. Litvin, «Categorical Analysis of Logical Networks in Application to Intelligent Radar Systems,» 2020 IEEE International Conference on Information and Communication Technologies. Science and Technology (PIC S&T), Kharkiv, Ukraine, 2020, pp. 235–238, DOI: 10.1109/PICST51311.2020.9467893
44. Cenzer, V. W. Marek and J. B. Remmel, «On the complexity of index sets for finite predicate logic programs that allow function symbols,» Journal of Logic and Computation, vol. 30, no. 1, pp. 107–156, Jan. 2020, DOI: 10.1093/logcom/exaa005
45. A. Kabulov, E. Urunbaev, and A. Ashurov, «Logical method for finding maximum common subsystems of Boolean equation systems,» 2020 International Conference on Information Sciences and Communication Technologies (ICISCT), Tashkent, Uzbekistan, 2020, pp. 1–5, DOI: 10.1109/ICISCT50599.2020.9351394

46. H. Qi, B. Li, R.-J. Jing, A. Proutiere and G. Shi, «Distributed Solving Boolean Equations over Networks,» 2020 59th IEEE Conference on Decision and Control (CDC), Jeju, Korea (South), 2020, pp. 560–565, DOI: 10.1109/CDC42340.2020.9304144
47. F. N. Castro, O. E. González and L. A. Medina, «Diophantine equations with binomial coefficients and perturbations of symmetric Boolean functions», in IEEE Transactions on Information Theory, vol. 64, no. 2, pp. 1347–1360, February 2018, DOI: 10.1109/TIT.2017.2750674
48. L. Gong, J. Zhang, L. Sang, H. Liu and Y. Wang, «The Unpredictability Analysis of Boolean Chaos,» in IEEE Transactions on Circuits and Systems II: Express Briefs, vol. 67, no. 10, pp. 1854–1858, Oct. 2020, DOI: 10.1109/TCSII.2019.2949571
49. T. Kosovska, «Implementation of a partial sequence of formulas for coarse-grained solution of AI problems within a logic-predicate approach,» 2019 Computer Science and Information Technologies (CSIT), Yerevan, Armenia, 2019, pp. 65–68, DOI: 10.1109/CSITechnol.2019.8895153
50. Y.-F. Tseng and S.-J. Gao, «Efficient encryption of predicate subsets for the Internet of Things,» 2021 IEEE Conference on Dependable and Secure Computing (DSC), Aizu-Wakamatsu, Fukushima, Japan, 2021, pp. 1–2, DOI: 10.1109/DSC49826.2021.9346245
51. K. Smelyakov, A. Chuprina, O. Bogomolov, and N. Gunko, «Effectiveness of neural network models for face detection and recognition,» 2021 IEEE Open Conference on Electrical, Electronic, and Information Sciences (eStream), 2021, pp. 1–7, DOI: 10.1109/eStream53087.2021.9431476
52. Sharonova N. Issues of fact-based information analysis [Electronic resource] / N. Sharonova, A. Doroshenko, O. Cherednichenko // Computational Linguistics and Intelligent Systems (COLINS 2018): Vol. 1: Lviv, 2018. P. 11–19. URL: <http://ceur-ws.org/Vol-2136/10000011.pdf>
53. D.E. Sitnikov, B. D’Kruz, P.E. Sitnikova, «Identification of essential data features by constructing and manipulating logical equations,» Data Mining II, WIT Press, 2000, 241–248.
54. Melnik K. Towards information technology for medical screening: an approach based on healthgrid / K. Melnik, O. Cherednichenko, V. Glushko // Information Systems: Methods, Models and Applications. – Heidelberg: Springer, 2013. P. 202–204.

DEEP REINFORCEMENT LEARNING-BASED ROBUST ROUTING FOR SCALABLE UAV SWARM COMMUNICATION

Sopov Ie., Artomova A.

Robust routing in mobile ad hoc and UAV swarming networks remains difficult due to rapid topology changes, interference, and intermittent links. We propose DeepCQ+, a scalable routing framework that augments the classical CQ/CQ+ family with a single multi-agent deep reinforcement learning (MADRL) policy trained under the centralized-training, decentralized-execution paradigm. The agent state uses a lightweight top-K neighbor encoding of CQ variables (confidence and hop estimates), while a PPO policy learns when to switch between unicast and adaptive flooding to balance delivery and overhead. To support and benchmark learning, we build an ns-3 (v3.39) slotted evaluation harness with RWP mobility and standard metrics (goodput; normalized MAC overhead). Baseline experiments (unicast only) over $N \in \{20, 30, 40, 50\}$ nodes and 4/8 flows show that mean goodput rises with density – approximately 0.66 ($N = 20$), 0.63 ($N = 30$), 0.72 ($N = 40$), 0.85 ($N = 50$) – while normalized overhead drops from ~ 0.15 ($N = 20 - 30$) to ~ 0.08 ($N = 40$) and ~ 0.062 ($N = 50$). These results quantify the operating regime where adaptive flooding is most valuable (notably $N = 30 - 40$, where variance is highest), and provide a reproducible baseline for learning-based control. The paper details the DeepCQ+ architecture, reward shaping for overhead awareness, and the ns-3/Matlab pipeline for data aggregation and plotting. Collectively, the study establishes a principled path toward routing policies that retain CQ+ robustness while scaling to larger swarms with lower control cost.

1. Introduction

Designing robust and efficient routing algorithms remains one of the most challenging problems in communication and computer networks. This challenge becomes even greater in tactical MANETs, where network conditions are highly dynamic and unpredictable due to node mobility, frequent topology changes, interference, and potential jamming [1]. Such factors substantially reduce reliability. Consequently, many traditional MANET routing protocols require frequent recomputation of end-to-end routes when detecting topology changes, which in turn causes periodic throughput losses while traffic is suspended during recomputation. These losses increase as the rate of topology change grows.

To mitigate these effects, broadcasting (or flooding), where a packet is transmitted to all neighbors, has become a popular approach for maintaining packet delivery performance in highly dynamic networks [2]. Classical MANET routing protocols such as OSPF and OLSR [3], perform reliably under stable conditions

but degrade significantly under rapid mobility. Adaptive extensions also fail to respond quickly enough in highly dynamic scenarios [4].

In this work, we focus on distributed routing algorithms that share only limited information – specifically, two floating-point values – via per-hop acknowledgment (ACK) packets. This approach is efficient since ACKs are already present in lower-layer protocols (e.g., MAC-layer acknowledgments) and require no additional implementation. The seminal work on Q-routing introduced reinforcement learning (RL), specifically Q-learning, to minimize delivery time by enabling each node to learn path qualities (Q-values) locally and update them through ACK exchanges [5]. Building on this, developed CQ-routing, which added confidence values (C-values) to improve adaptability in dynamic networks. To further enhance reliability, the Smart Robust Routing (SRR) algorithm [6] introduced selective broadcasting, enabling nodes to decide when to broadcast based on heuristic rules. This extension, referred to as CQ+ (or Robust Routing for Dynamic Networks, R2DN), effectively balances reliability and overhead. However, CQ+ relies on a single network parameter – the best-path confidence level – to switch between unicast and broadcast. This limited perspective often results in locally optimal decisions that fail to fully account for rapid topology changes or congestion levels. While CQ+ consistently delivers high packet delivery ratios across many scenarios, it incurs a significant broadcast overhead, making it an important but imperfect baseline for more advanced frameworks.

This motivates the question: can emerging deep reinforcement learning frameworks reduce overhead while maintaining, or even improving, CQ-routing performance (particularly goodput)? Routing decisions such as next-hop selection are natural candidates for RL-based techniques. Since Boyan’s original Q-routing protocol [7, 9], a range of RL-inspired methods have been proposed for packet routing and scheduling. However, these approaches typically scale poorly with increasing network size or when system parameters change. To address these limitations, recent studies have explored multi-agent deep reinforcement learning (MADRL) [8], though many of these designs remain constrained by scalability issues, as they require retraining whenever new nodes are added.

To the best of our knowledge, no prior study has simultaneously achieved both scalability and robustness using MADRL in dynamic networks such as MANETs. Our proposed DeepCQ+ overcomes these limitations by embedding MADRL into the CQ+ framework. Specifically, MADRL agents control next-hop flooding decisions while retaining the structural simplicity of CQ+. This design enables a single trained agent to produce scalable, robust routing policies applicable to

any node, regardless of network size. More importantly, the DeepCQ+ methodology allows training under limited scenarios – small networks, selective traffic patterns, or constrained mobility – while remaining effective when deployed in environments outside the training distribution. This property addresses the curse of dimensionality and mitigates catastrophic forgetting, making DeepCQ+ a significant advancement toward reliable routing in highly dynamic networks.

2. Resilient Routing Architecture

The SRR algorithm extends the classical CQ-routing by incorporating two key network parameters: the H -factor and the C -value [9, 10].

For each node i , the H -factor, denoted as $h(i, j, d)$, represents an estimate of the minimum number of hops from node i to destination d via a potential next-hop j . To capture network dynamics, each node also maintains a confidence level, or C -value $c(i, j, d)$, which reflects the likelihood that a packet will successfully reach destination d through j .

The C -value increases with successful transmissions (acknowledgements received) and decays with each transmission attempt. Both C - and H -values are updated using c_{ack} and h_{ack} parameters, which are propagated through ACK packets from the receiving node back to the transmitter. At the next-hop node j , the ACK-propagated update terms are defined as in (1) and (2):

$$h_{ack} = 1 + h(j, \hat{k}, d) \quad (1)$$

$$c_{ack} = c(j, \hat{k}, d) \quad (2)$$

where $\hat{k} = \arg \min_k h(j, k, d)(1 - c(j, k, d))$ is the optimal next-hop selected by node j for reaching destination d .

Thus, the C - and H -values act as exponential moving averages of c_{ack} and h_{ack} . In case of a failed transmission (i.e., no ACK received), the H -value cannot be updated, while the C -value is degraded by setting $c_{ack} = 0$, reflecting path failure. The recursive updates are given in (3) – (4):

$$h_{t+1}(i, j, d) = (1 - \alpha)h_t(i, j, d) + \alpha h_{ack} \quad (3)$$

$$c_{t+1}(i, j, d) = \begin{cases} (1 - \lambda)c_t(i, j, d), & \text{failure} \\ (1 - \lambda)c_t(i, j, d) + \lambda c_{ack}, & \text{other wise} \end{cases} \quad (4)$$

where $0 \leq \alpha \leq 1$ is the discount factor for new observations (adaptively set as $\alpha = \max(c_{ack}, 1 - c_t(i, j, d))$) and λ is the decay parameter controlling the update rate. When a packet reaches destination d , both c_{ack} and h_{ack} are set to 1, reflecting full confidence at one hop away.

The CQ+ routing protocol consists of three main steps:

1. Reception of ACK packets and update of C - and H -values;
2. Reception of data packets, with duplicate detection, queuing, or delivery if the node is the destination;
3. Transmission of data packets via either unicast or broadcast according to the routing policy.

The routing policy selects the next-hop that minimizes the joint metric of path uncertainty and expected hops (5):

$$j^* = \arg \min_j h(i, j, d)(1 - c(i, j, d)) \quad (5)$$

If sufficient confidence exists in the chosen next-hop, the packet is unicasted. Otherwise, the protocol probabilistically triggers broadcast to reduce uncertainty, with probability defined as:

$$P_{BC} = \varepsilon + (1 - c(i, j^*, d))(1 - \varepsilon) \quad (6)$$

where ε is a small exploration parameter ensuring a minimum probability of broadcast.

In the enhanced framework, the broadcast/unicast decision of the CQ+ protocol is replaced by a multi-agent deep reinforcement learning (MADRL) model, enabling more robust and adaptive routing in highly dynamic networks.

3. CQ Routing with Multi-Agent Deep Reinforcement Learning (DeepCQ+)

Our robust routing design follows the decentralized partially observable Markov decision process (Dec-POMDP), a standard framework for cooperative multi-agent decision-making under uncertainty [14]. In this setting, each agent observes only local information and optimizes its behavior while accounting for environmental and inter-agent uncertainty.

Policy optimization. We employ policy gradient (PG) methods that optimize the policy directly by ascending the expected discounted return [11]. Specifically, we adopt Proximal Policy Optimization (PPO) [16], which performs multiple epochs of stochastic gradient ascent per update and is known for its stability and ease of implementation.

Centralized training, decentralized execution. Due to partial observability and communication constraints, agents must execute decentralized policies based solely on local observations. Nevertheless, these policies can be trained centrally in simulation, a paradigm widely used in MARL [12].

Parameter sharing. For homogeneous agents, we apply parameter sharing so that all nodes use a single policy π_θ , improving scalability and convergence [13]. Fig. 1 illustrates the multi-agent routing environment, centralized training with decentralized execution, and shared policy parameters.

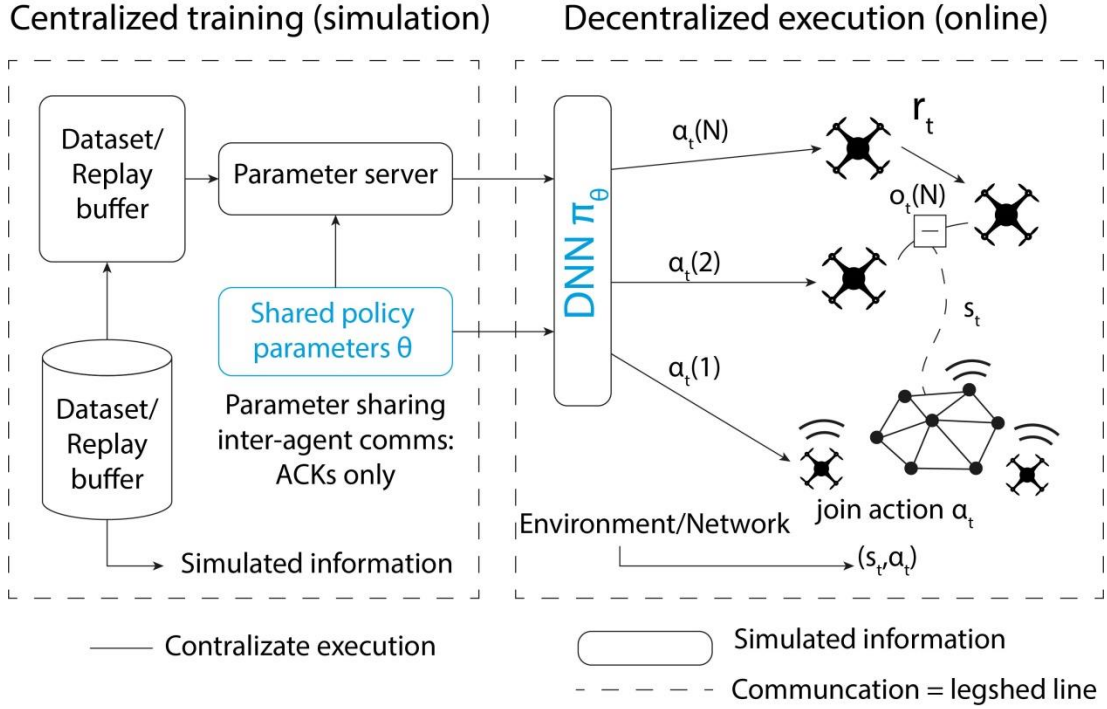


Fig. 1. Multi-agent routing environment with shared policy parameters

Training is centralized and execution is decentralized. At time t , each agent i applies the common policy π_θ to its local observation $o_t(i)$ and selects an action $a_t(i)$. The environment evolves according to the joint action of all agents, producing the next state s_{t+1} and per-agent rewards.

4. DeepCQ+ Design Framework and Approach

We consider a homogeneous MANET with variable network size N (e.g., $5 \leq N \leq 50$) and multiple unicast flows with randomized source-destination pairs. Nodes move with random velocities/directions; baseline mobility follows Gauss – Markov and random waypoint models, further diversified as discussed in

Section IV-A. Each RL episode spans at least T time slots (packet-duration slots). The episode terminates once traffic exceeds $T_{\max}$ and queues drain; no new traffic arrives after $T_{\max}$. Goodput is the ratio of successfully delivered packets to injected packets (duplicates at the destination excluded). A single shared policy π_θ (one DNN for all nodes) is trained on a narrow distribution (single network size/flow, modest dynamics) and tested across wider ranges (larger N , multiple flows, higher mobility).

We retain only the best K neighbors (e.g., $K = 4$) per node based on their (C, H) values. For notational simplicity, we drop the explicit destination index in c_t and h_t . The best- K neighbor vectors are defined as in (7):

$$c_t(i) = [c_t(i, i_1), c_t(i, i_2), \dots, c_t(i, i_K)], \quad h_t(i) = [h_t(i, i_1), h_t(i, i_2), \dots, h_t(i, i_K)] \quad (7)$$

where neighbors $i, \dots, i_K$ are ordered by the CQ+ metric as in (8),

$$h_t(i, i_1)(1 - c_t(i, i_1)) \leq \dots \leq h_t(i, i_K)(1 - c_t(i, i_K)) \quad (8)$$

so that the most promising next hops appear first.

State/observations. To capture temporal dynamics, we feed a fully connected neural network (FCNN) with first-order differences of observations. The policy input at node iii is defined in (9):

$$o_t(i) = [c_t(i), h_t(i), \Delta c_t(i), \Delta h_t(i), a_{t-1}(i), p_{t-1}(i)] \quad (9)$$

where $\Delta c_t(i) = c_t(i) - c_{t-1}(i)$, $\Delta h_t(i) = h_t(i) - h_{t-1}(i)$; $a_{t-1}(i)$ is the previous action at node i ; and $p_{t-1}(i)$ indicates whether the current packet arrived via broadcast or unicast.

Tabl. 1 summarizes capabilities across CQ-routing [7], SRR (CQ+) [8], and DeepCQ+ (this work): CQ values, confidence/broadcast, and MADRL control.

Table 1

**Comparative Analysis of CQ-based Protocols vs. DeepCQ+
in Dynamic MANETs**

Routing Protocol	Control logic	Scalability	Overhead	Goodput	Notes
CQ-routing [7]	Heuristic (Q-values)	Medium (node-local)	Low (unicast only)	Medium; degrades at high mobility	Baseline RL-inspired
SRR (CQ+) [8]	Heuristic + selective flooding	Medium	High (uncertainty-driven flooding)	High; robust under dynamics	Uses C/H + broadcast policy
DeepCQ+ (this work)	Learned (MADRL)	High (CTDE + sharing)	Medium ($\downarrow$ vs CQ+)	High ($\geq$ CQ+)	Adaptive unicast/broadcast

CQ-routing is a heuristic baseline that uses Q-values only; SRR (CQ+) augments CQ with selective flooding, improving robustness (Goodput) at the cost of higher Overhead. DeepCQ+ replaces the broadcast/unicast switch with a learned MADRL policy (CTDE + parameter sharing), preserving high Goodput while reducing Overhead versus SRR and scaling better with network size.

5. The DeepCQ+ Reward Function

We first align with CQ+ by defining a stochastic routing policy $\pi(a|s_t)$ and rewards for broadcast/unicast decisions as in (10):

$$r_t(s_t, a_t) = \begin{cases} 1 - c_t(i, i_1)\tilde{\varepsilon}, & a_t = 1 \text{ (broadcast)}, \\ c_t(i, i_1)\tilde{\varepsilon}, & a_t = 0 \text{ (unicast)}, \end{cases} \quad (10)$$

where $\tilde{\varepsilon} = 1 - \varepsilon$. This mirrors CQ+ behavior for small exploration ε , providing a consistent baseline.

To explicitly discourage overhead, we define overhead as the ratio of total transmissions N_{TX} to delivered packets N_D , normalized by network size N : $OH = \frac{1}{N} \cdot \frac{N_{TX}}{N_D}$. We also distinguish excess transmitted bits per delivered bits and (ii) total (data+ACK) bits per delivered data bit. Incorporating delivery, failure, and ACK-induced replication, the overhead-aware reward is given in (11):

$$r_t = w_1 1_D - w_2 1_Z - w_3 \frac{N_{ack}}{N} \quad (11)$$

where 1_D rewards successful delivery when node i contributes; 1_Z penalizes transmissions without any received ACKs; and N_{ack} approximates packet replication/overhead. Weights w_1, w_2, w_3 are tuned to minimize overhead while preserving delivery. This tuned reward defines the DeepCQ+ routing policy, whose algorithms are summarized in Algorithm 1. The proposed DeepCQ+ routing (CTDE with parameter sharing):

Inputs: size range N ; horizon T ; $T_{\max}$; neighbor cap K (e.g., $K = 4$); exploration ε ; reward weights w_1, w_2, w_3 ; PPO hyperparameters; initial C/H tables. Outputs: shared policy π_θ ; decentralized routing decisions.

Phase I – Centralized training (simulation)

1. Initialize policy π_θ and value estimator; initialize C/H tables.
2. For each episode: randomize N , flows, mobility (Sec. IV-A); reset queues and C/H .

3. For $t = 1, \dots, T$ until traffic $> T_{\max}$ and queues are empty:
 - 3.1 For each node i : build top- K neighbor vectors $c_t(i), h_t(i)$ as in (7); order neighbors by CQ+ metric as in (8).
 - 3.2 Form observation $o_t(i)$ with first-order differences and indicators as in (9).
 - 3.3 Sample action $a_t(i) \sim \pi_\theta(o_t(i))$
 - 3.4 If unicast: select j^* by (5) and transmit; else broadcast per policy (CQ+ context (6)).
 - 3.5 Process ACKs; compute c_{ack}, h_{ack} by (1)–(2), update C, H via (3)–(4).
 - 3.6 Compute reward using (10) or (11); store o_t, a_t, r_t, o_{t+1} .
4. Run PPO updates (policy gradient) on collected trajectories; repeat episodes until convergence.

Phase II – Decentralized execution (deployment)

5. Freeze π_θ ; deploy the same π_θ at all nodes (parameter sharing).
6. For each arriving packet at node i : build top- K per (7) – (8); form $o_t(i)$ (9); sample $a_t(i)$.
7. Forward: unicast via j^* (5) or broadcast; update C, H with (1) – (4); suppress duplicates; continue until delivered.
8. Track Goodput and Overhead online (no gradient updates during execution).

We propose DeepCQ+, a multi-agent deep reinforcement learning (MADRL) extension of CQ routing with centralized training, decentralized execution, and shared parameters across nodes. The agent state uses top- K neighbors and lightweight temporal features; the learned policy adaptively switches between unicast and broadcast, while ACK-based updates ensure robustness to rapid topology changes. PPO training provides scalability and stability. Overall, DeepCQ+ achieves higher goodput with lower overhead than SRR (CQ+).

6. Experimental Setup and Evaluation

All experiments were carried out in ns-3 v3.39 using the ad-hoc Wi-Fi stack (Yans PHY, Adhoc MAC) and an event-driven slotted simulation runner. The deployment region is a square of side $L = 300$ m. Terminals (UAV/mobile nodes) follow the Random Waypoint mobility model with pause time drawn from the default ns-3 setting and speed $v \in [1, 2] \text{ ms}^{-1}$, re-sampled at waypoints. Each episode

comprises $T = 3000$ time slots of duration 20 ms; traffic injection ceases after $T_{\max} = 2000$ slots and the simulation continues until all queues drain.

We evaluate network sizes $N \in \{20, 30, 40, 50\}$ under 4 and 8 simultaneous unicast flows (randomized source–destination pairs per episode). Unless otherwise stated, figures report means across the two traffic intensities with 95% confidence intervals computed from the sample variance across intensities (normal approximation; see Statistical note below).

Two performance indicators are used:

- Goodput: ratio of successfully delivered data packets to injected packets (duplicates at the destination excluded);
- Normalized overhead (OH2): total MAC transmissions (data + ACK + control) divided by the number of delivered data packets and normalized by N (lower is better).

Aggregated numerical outcomes are reported in Table 2; trends are visualized in Figs. 1–3.

Table 2

Baseline MANET performance under RWP mobility

N	tag	N	flows	Injected	Delivered	Goodput	OH1	OH2
1	baseline	20	4	8000	5182	0.64775	0.142667	0.142667
2	baseline	20	8	16000	10697	0.668652	0.144475	0.144475
3	baseline	30	4	8000	5825	0.728125	0.125345	0.125345
4	baseline	30	8	16000	8536	0.533500	0.175832	0.175832
5	baseline	40	4	8000	6517	0.814625	0.0671513	0.0671513
6	baseline	40	8	16000	10067	0.629188	0.0911543	0.0911543
7	baseline	50	4	8000	6904	0.863000	0.0610370	0.0610370
8	baseline	50	8	16000	13555	0.847187	0.0621173	0.0621173

For sparse to moderate densities ($N = 20 - 30$), goodput is 0.53–0.73 with normalized overhead $\text{OH2} \approx 0.12 - 0.18$. For denser networks ($N = 40 - 50$), goodput rises to 0.63–0.86 while OH2 drops to 0.06 – 0.09. The spread between 4 and 8 flows explains the error bars shown in the figures.

As shown in Fig. 2 (Goodput vs. N , 95% CI), performance remains essentially flat at low densities and increases markedly with N : mean goodput ≈ 0.66 for $N = 20$, ≈ 0.63 for $N = 30$, ≈ 0.72 for $N = 40$, and ≈ 0.85 for $N = 50$. The broader confidence bands at $N = 30 - 40$ arise from sensitivity to load (4 vs. 8 flows): under higher offered load, transient queueing and route changes depress delivery;

adding nodes at $N = 50$ provides alternative paths and restores end-to-end throughput. These trends are consistent with the raw outcomes in Tab. 2.

Fig. 3 shows a monotonic decrease in OH2 beyond $N = 30$: from ≈ 0.15 at $N = 20 - 30$ down to ≈ 0.08 at $N = 40$ and ≈ 0.062 at $N = 50$. This reduction is attributable to shorter average hop counts and fewer retransmissions as topology becomes richer, which offsets increased contention. The larger variance at $N = 30$ mirrors the goodput spread and indicates a regime where mobility and flow density interact unfavorably for the baseline policy.

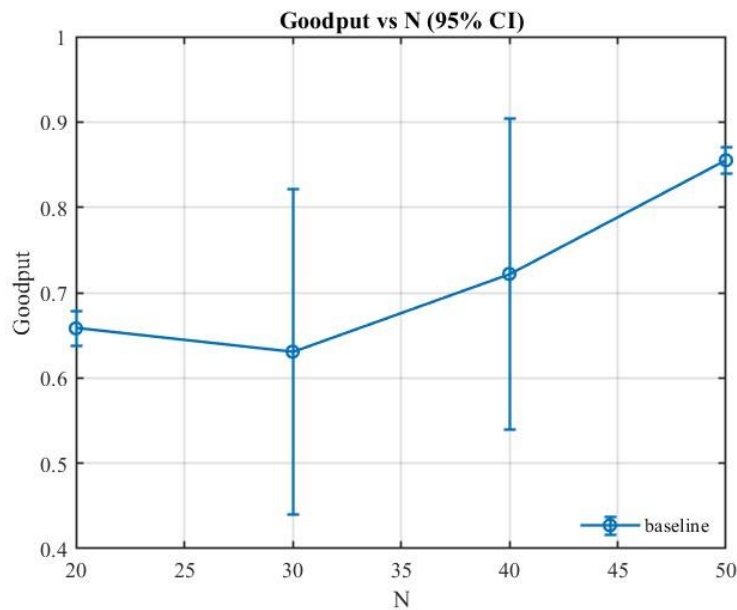


Fig. 2. Impact of Node Count on Goodput for Unicast MANET Routing (95% CI)

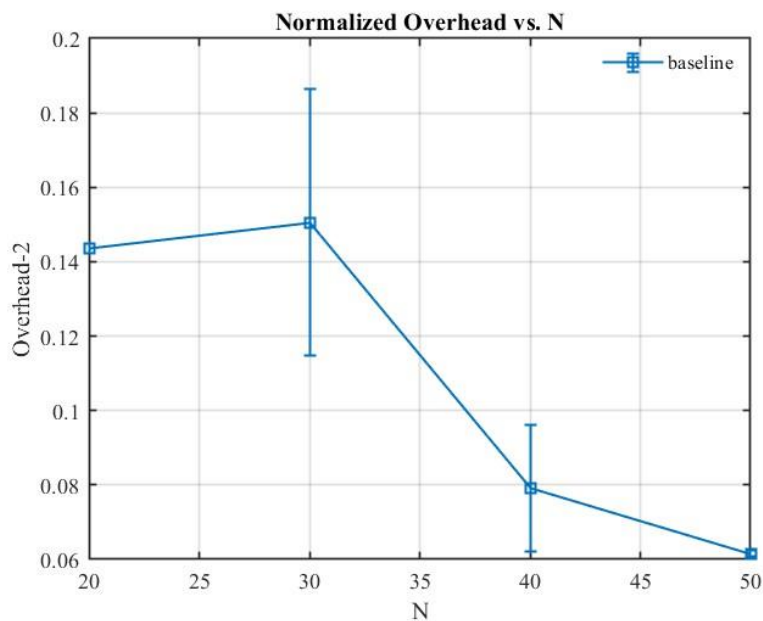


Fig. 3. Normalized Overhead (OH2) vs. Network Size under RWP Mobility

For completeness, Fig. 4 reproduces goodput vs. N with alternate y-axis scaling for direct comparison against CQ+ and DeepCQ+ in the next section.

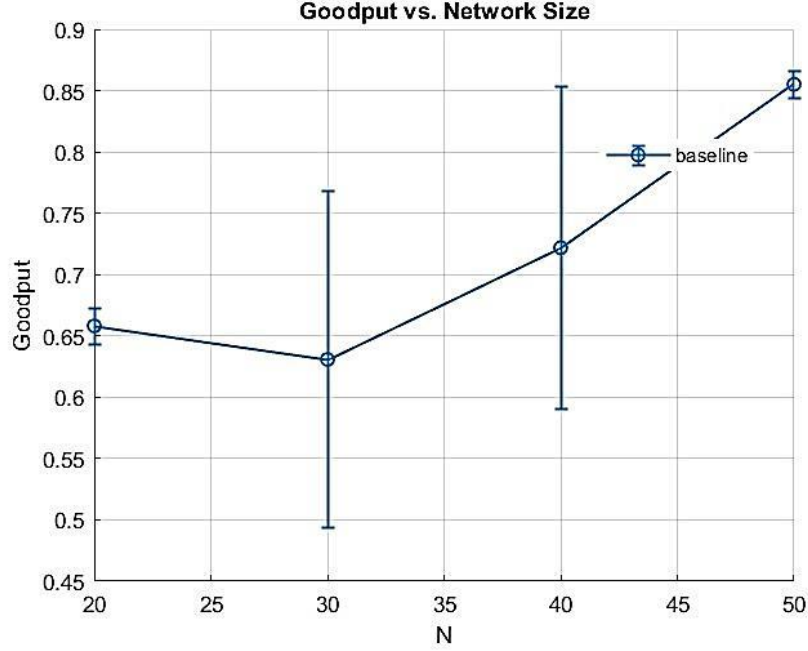


Fig. 4. Baseline Goodput across N (y -axis harmonized for comparison)

Even without selective flooding or learned control, increasing node density in this mobility regime simultaneously improves goodput and reduces normalized overhead. These baselines (Table 2, Figs. 1–3) constitute the quantitative reference against which gains of CQ+ and DeepCQ+ will be assessed (e.g., maintaining high goodput already at $N = 20 - 30$ while avoiding the overhead levels observed in Fig. 2).

Per N , we aggregate two operating points (4 and 8 flows). The mean is reported across these points; error bars in Figs. 1–3 correspond to 95% confidence intervals computed as $1.95\sigma/\sqrt{n}$ with $n = 2$ (normal approximation). With $n = 2$, CIs are indicative rather than definitive; in subsequent sections we increase the number of runs per setting when comparing CQ+/DeepCQ+.

7. Conclusions

This work introduced DeepCQ+, a scalable and robust routing framework that augments the CQ+/SRR structure with a single shared multi-agent policy learned under the CTDE paradigm. By retaining CQ+’s lightweight per-hop state while delegating the unicast/broadcast switch to a learned controller, DeepCQ+ targets the dual goals that prior MADRL routing attempts have struggled to achieve

simultaneously: scalability across network sizes and robustness under mobility and load variation.

In preparation for the learning-based controller, we built and validated an ns-3 evaluation harness with slotted execution, RWP mobility, and standard metrics. Baseline (unicast-only) experiments show two consistent trends across $N \in \{20, 30, 40, 50\}$ goodput rises with node density ($\approx 0.66, 0.63, 0.72, 0.85$, respectively), and normalized overhead (OH2) falls monotonically beyond $N = 30$ ($\approx 0.15 \rightarrow 0.08 \rightarrow 0.062$), reflecting shorter routes and fewer retransmissions in denser topologies. These effects persist across 4 and 8 flows, with larger variance at $N = 30 - 40$ due to load sensitivity. Taken together, the baselines set a quantitative reference against which CQ+ and DeepCQ+ will be assessed in the sequel.

The monotonic OH2 reduction indicates that broadcast need not be frequent in moderately dense swarms; conversely, the variance band at $N = 30 - 40$ reveals the regime where adaptive flooding is most valuable. DeepCQ+ is expressly designed to exploit these regimes, learning when to broadcast to preserve goodput at low/medium N while avoiding the overhead spikes seen in the heuristic CQ+ switch.

Current results use RWP mobility, a single physical/MAC stack, fixed packet duration, and only two traffic intensities per N . Confidence intervals are based on $n = 2$ operating points and are therefore indicative rather than definitive; subsequent comparisons (CQ+ vs. DeepCQ+) will increase repetitions per setting and diversify mobility (e.g., Gauss–Markov), traffic, and channel models.

We will integrate the learned broadcast/unicast policy into ns-3 and compare against CQ+ under identical scenarios; perform ablations on the top- K neighbor encoding and reward weights; evaluate generalization to untrained N , flow counts, and mobility/jamming patterns; and report energy and fairness metrics in addition to goodput/overhead. These steps align with the monograph structure and presentation rules (tables/figures and conclusions), ensuring compliance with authoring requirements.

References

1. Machakaire t. Enhancing manets for military applications: a comprehensive review of innovations, challenges, and research gaps // i-Manager's Journal on Wireless Communication Networks – T. 13. – №. 2. – 2025. DOI: <https://doi.org/10.26634/jwcn.13.2.22063>
2. Ismatov, A., Suh, B. K., Kim, J., Park, Y. B., & Kim, K. I. Adaptive Broadcast Scheme with Fuzzy Logic and Reinforcement Learning Dynamic Membership Functions in Mobile Ad Hoc Networks. *Mathematics*, 13(15), 2367. 2025. DOI: <https://doi.org/10.3390/math13152367>
3. Ostojic, Zivka Slepcevic. Comparative Studies of Link-State Routing Protocols in Wired IP and Ad-Hoc Wireless Network: OSPF, IS-IS and OLSR. Diss. University of Applied Sciences Technikum Wien, 2022.

4. Harib, Mouhcine, Hicham Chaoui, and Suruz Miah. «Evolution of adaptive learning for nonlinear dynamic systems: A systematic survey.» *Intelligence & Robotics* 2.1 (2022): 37-71.
5. Kunzel, Gustavo. «Application of reinforcement learning with Q-learning for the routing in industrial wireless sensors networks». (2021).
6. Darabkh, K. A., Al-Khazaleh, H. F., Al-Zubi, R. T., Alnabelsi, S. H., & Salameh, H. B. Efficient routing protocol for optimal route selection in cognitive radio networks over IoT environment. *Wireless Personal Communications*, 129(1), 209–253. (2023). DOI: <https://doi.org/10.1007/s11277-022-10093-6>
7. Martins, M. S., Sousa, J. M., & Vieira, S. A systematic review on reinforcement learning for industrial combinatorial optimization problems. *Applied Sciences*, 15(3), 1211. (2025). DOI: <https://doi.org/10.3390/app15031211/>
8. Gronauer, S., Diepold, K. Multi-agent deep reinforcement learning: a survey. *Artif Intell Rev* 55, 895–943 (2022). DOI: <https://doi.org/10.1007/s10462-021-09996-w/>
9. R. E. Ali, B. Erman, E. Bas, tug, and B. Cilli, «Hierarchical deep double Q-routing», in ICC 2020-2020 IEEE International Conference on Communications (ICC). IEEE, 2020, pp. 1–7.
10. Kaviani, S., Ryu, B., Ahmed, E., Larson, K., Le, A., Yahja, A., & Kim, J. H. DeepCQ+: Robust and scalable routing with multi-agent deep reinforcement learning for highly dynamic networks. In MILCOM 2021-2021 IEEE Military Communications Conference (MILCOM) (pp. 31–36). IEEE. (2021, November).
11. Prashanth, L. A., & Fu, M. C.. Risk-sensitive reinforcement learning via policy gradient search. *Foundations and trends® in machine learning*, 15(5), 537–693. (2022)
12. Zhang, K., Yang, Z., & Başar, T. Multi-agent reinforcement learning: A selective overview of theories and algorithms. *Handbook of reinforcement learning and control*, 321–384. (2021).
13. Hu, F., Deng, Y., & Aghvami, A. H. Scalable multi-agent reinforcement learning for dynamic coordinated multipoint clustering. *IEEE Transactions on Communications*, 71(1), 101–114. (2022).

BEST PRACTICES FOR SYSTEMATIZING ROUTINE PROCESSES IN MOBILE DEVELOPMENT PROJECTS

Troshchyi D., Kuznetsova Yu.

In mobile development, a significant amount of time is spent on repetitive tasks: setting up the environment, configuring the build, checking pull requests, updating dependencies, running tests, etc. Despite the availability of automation tools such as CI/CD, templates, and scripts, their implementation is often complicated by technical, organizational, and mental barriers. Typical difficulties include unstable environments, fragmented solutions, lack of support, and the absence of unified approaches. Automation is often implemented as a one-off initiative without a long-term vision, which reduces its viability. An analysis of recent studies shows that most of them focus on testing, without covering other critical stages of the application lifecycle. Existing research rarely takes into account the specifics of mobile platforms, such as the presence of markets, signatures, and SDK updates. As a result, teams create their own solutions that are difficult to scale or quickly become obsolete. The article summarizes the problems that developers face at all stages, from project configuration to maintenance. It has been established that routine tasks affect not only the pace of work and quality, but also team motivation, contributing to burnout and the accumulation of technical debt. Therefore, we propose creating a framework that systematically links typical tasks with tools, practices, and viability criteria. This approach will reduce fragmentation and increase the stability and efficiency of mobile development.

Introduction

In mobile app development, a significant part of daily work consists of repetitive tasks that may seem minor at first glance. These include environment and build configuration [1], request pool checking [2], building, dependency updates [3], test running [4], and so on.

Such routine tasks accompany virtually every stage of an application's life cycle, from design to maintenance, and they often take up a significant portion of a developer's working time.

The next step after identifying such tasks is to try to improve their efficiency through automation [5], templating [6], or organizational changes. However, in practice, even the most obvious and simple improvements are often hindered by difficulties such as lack of time, technical limitations, complexity of maintenance, or simply the absence of unified approaches within the team. As a result, some repetitive tasks are not solved systematically, but are done manually each time, despite the availability of tools or experience from previous projects.

The question is not only how to perform a routine task efficiently, but also why attempts to solve it often encounter resistance: technical, organizational, or mental.

In mobile development practice, such bottlenecks only deepen over time, especially in the context of accelerated release cycles, constant platform updates, and an increasing number of dependencies.

1. Complexity of solving routine tasks

At first glance, most repetitive tasks in mobile development seem pretty basic: set up a build, check a variable in CI (Continuous Integration), handle an API (Application Programming Interface) error, write a basic unit test. These are things that don't seem to need any fancy solutions. However, this is precisely where the problem lies – attempts to automate or standardize them often encounter hidden barriers that, when combined, create a significant obstacle.

First, there are technical limitations. Even modern platforms – Android and iOS – have their own peculiarities that complicate the unification of approaches. For example, updating libraries in an Android project may seem simple, but in a project with a large number of modules and no centralized location for dependencies, this turns into a series of manual checks and edits in several configuration files. In projects with complex CI/CD pipelines, adding new steps, even formally simple ones, sometimes requires diving into Jenkins or Bitrise documentation, and sometimes simply writing a separate Bash script.

Second, organizational barriers. In many teams, attempts to systematize or automate routines are not supported by management. The reasons are trivial: tight deadlines, a focus on visual functionality rather than processes, or a lack of internal culture of fixing what works. In such conditions, even obvious improvements, such as introducing a template generator for creating new screens, are put off until later, which never comes.

Third, the human factor. Many developers find it easier to do a task themselves than to spend time on «automation for automation's sake». This is partly justified; indeed, not every routine task requires a technical solution. But when there are dozens of such actions every day, their total weight becomes significant. This is especially true if these tasks are performed by different people with varying levels of quality, which creates problems in terms of support, documentation, and knowledge transfer.

Another common problem is technical debt associated with previous attempts at a solution. For example, a team may have implemented a tool to automate assembly, but over time it stopped working due to API changes or platform limitations, and no one took on the task of maintaining it. As a result, the tool

remained in the project but is not used. New developers are forced to look for workarounds, which are again not automated.

Equally important is that mobile development is often ignored in general automation practices. Many of the existing solutions and recommendations were developed in the context of web or backend development, where the infrastructure is more predictable: deployment happens without involving app stores, new changes can be uploaded to the server at any time, and CI/CD is often already integrated into the platform. In mobile development, the situation is different – SDK (Software Development Kit) updates, strict App Store and Google Play requirements, release signing – all of this create an additional workload that is difficult to automate without deep knowledge of the platform.

So, the problem of solving «routine» tasks in mobile development is not just the lack of technical tools. Often, the solution already exists, but it doesn't take root because it requires regular support, team agreement, and flexibility in changing processes. In addition, there is a lack of a common approach to automation and efficiency improvement: teams create their own templates, scripts, and utilities, but these solutions remain local, difficult to scale, and transferable to other projects. As a result, even with the technical capabilities in place, a significant portion of repetitive tasks continue to be performed manually, with all the associated time losses.

2. Overview of typical difficulties in research

Despite the widespread popularity of mobile development in practice, the scientific base devoted to systematic analysis and routine problem solving in this context remains rather limited. Most publications either focus on individual aspects of the mobile application lifecycle, such as testing, or are of a general nature that is not adapted to the specifics of mobile engineering. At the same time, there are virtually no publications that comprehensively cover all phases of development, from project configuration and requirements engineering to release stages, monitoring, and feedback processing.

In their work «A survey on the practices of mobile application testing» [7], the authors conduct a large-scale survey of mobile developers on their approaches to testing. The study confirms that manual testing prevails, while automated tools are implemented fragmentarily, often without proper integration into the overall development infrastructure. The low level of unit test coverage, lack of UI gesture support, and disregard for scenarios with contextual interaction

(geolocation, camera, push notifications) are highlighted separately, turning many checks into repetitive routines.

A broader overview is provided in «A Survey on Mobile App Development Approaches with the Industry Perspective» [8], which examines not only technical approaches but also organizational development models. The authors point out that mobile engineering is often based on Agile or MVP cycles, but there are almost no established models that adapt classic software engineering principles to the specifics of mobile development. In particular, it is noted that handling changes in requirements, adapting to platform updates, and rapid prototyping remain predominantly manual processes, with automation affecting only a small portion of the life cycle, confirming the key thesis about the limitations of solutions for «routine» tasks.

The paper «The Applicability of Automated Testing Frameworks for Mobile Application Testing» [9] reviews 56 scientific papers analyzing the capabilities of existing frameworks for mobile testing. The study is systematic and covers the evolution of methods from keyword-driven to hybrid and AI-assisted approaches. However, even the most modern tools demonstrate limited adaptation to mobile conditions: poor support for complex scenarios, high configuration requirements, and a lack of flexibility to support different platform versions. These factors significantly limit the effective resolution of routine tasks during the testing phase, especially when it comes to checking UI components and regressions.

The study «How do Developers Test Android Applications» [10] confirms that even in Android projects with stable teams, automation is used to a limited extent. The authors emphasize that developers often consider testing to be either secondary or not worth the cost. This creates a culture of manual functionality review, where stability is ensured not by tools, but by the personal time and effort of developers. Ultimately, this model leads to repetitive actions that remain in the project for years.

Overall, the sources analyzed show that even where there are attempts to solve routine tasks, such as testing, these solutions are highly specialized, often not scalable, and require significant human resources to maintain. The most systematic of the existing approaches focus exclusively on testing but hardly touch on other critically important stages of the life cycle. The lack of cross-stage integration and methodological consistency confirms the need for a framework that will combine routine tasks with the appropriate tools and practices within a single structure adapted to mobile development.

3. Problems of repetitive actions at different stages of development

Mobile app development consists of repetitive actions that occur at every stage: from the first launch of the IDE (Integrated Development Environment) to technical support after release. Although at first glance each of these tasks seems fairly simple on its own, together they create a significant workload that is not always easy to reduce, even with the right tools. The reason is not only technical limitations, but also the fact that most solutions are created on an ad hoc basis, without a systematic approach, long-term vision, or responsibility for support.

Problems begin as early as the configuration stage of a new project. Even within a single platform (e.g., Android), the structure of an application can vary from team to team. Some use modular architecture, while others use a monolithic one. The same applies to build schemes, signing config, integration with CI/CD (Continuous Integration / Continuous Delivery / Deployment), Firebase, analytics, libraries for working with HTTP (HyperText Transfer Protocol), DI (Dependency Injection), and so on. As a result, instead of reusing typical configurations, developers are forced to reconfigure everything from scratch each time, copying fragments from previous projects or documentation. In the best case, internal boilerplate is used; in the worst case, each engineer creates a structure based on their own experience.

Any attempt to standardize these actions often runs into a lack of time, support, or unwillingness to change what works. Writing the code itself also involves repetitive patterns and complexities that are difficult to automate.

In addition to generating basic elements (ViewModel, layout, models), there is a lot of manual work involved in creating the UI. If a ready-made component library is not used, most of the interface is laid out manually, either from scratch or based on Figma layouts. Moreover, adapting to different screen sizes (especially for Android) is still a challenging task. Automatic UI tests do not always guarantee correct rendering, and visual snapshot tests, although popular, often require manual approval and are too sensitive to minor changes.

Another common problem is code duplication. Its appearance is a direct consequence of the lack of centralized solutions. For example, if the team does not have shared components for navigation, error handling, logging, or status display (e.g., loading/error/empty), each screen is implemented in its own way. Over time, dozens of classes with similar logic appear, which are difficult to update simultaneously or extract into a common layer. Ways to reduce duplication, such as templates or wrappers, are often not implemented due to lack of time, and repetitive code accumulates.

The process of writing code itself is also part of the «routine» that people try to speed up. For example, with autocomplete, snippets, internal plugins, or generators. But supporting such solutions in complex projects requires constant adaptation to changes. Often, a team starts out enthusiastically, creating a set of utilities or templates, but after a few months, they stop working with new versions of tools or don't support changed requirements. Without a defined update process, these developments are simply forgotten, and routine work returns.

CI/CD is an area that is formally best suited for automation, but in practice it is also vulnerable. Any update to SDK or Google Play requirements can cause failures. The complexity lies in the fact that the build process is a combination of scripts, keys, flavors, signatures, environment variables, and saved artifacts.

Often, CI/CD configuration depends on the person who created it, and if that person leaves the team, adapting or scaling pipelines becomes a problem. In such cases, instead of automation, the project reverts to manual build delivery because it is easier in the short term. Working with tests is no exception.

Ideally, tests should be part of the development process, but when the application structure changes, old tests start to break. If there is no clear responsibility for maintaining the test environment, or if the tests themselves were created without following best practices, then any large-scale update will cause them to fail en masse, and they will simply be disabled. This also applies to UI tests, which often depend on the stability of screens, locales, backend connections, animation timings, and so on.

The release and regression stages remain a serious source of routine work. Preparation for publication should formally be as automated as possible, but in reality it is often accompanied by manual verification, copying change descriptions, and preparing screenshots and specifications for different app stores. Even minor actions, such as updating a version, publishing a release branch, or generating a changelog, are performed manually in many projects, which wastes time and creates risks of errors. The reason is that automation is not fully implemented or is outdated. As a result, when the release deadline approaches, the team simply repeats the same actions they have already performed dozens of times, without any changes.

A separate challenge is post-release regression, when new functionality unexpectedly starts to break. This often happens not because something is fundamentally broken, but because the change was not sufficiently tested under normal conditions. Testing may have been only local, or performed on a different build flavor, or someone manually fixed something in production during the build. In this case, the search for causes begins: is it a bug in the code, in the configuration, in CI, in Google Play, or in the tester's environment? This is a classic situation

where the team spends a day or two localizing a bug that arose due to an obscure deviation in the settings or library version.

Even worse is when complaints from users appear after release, and the team only learns about the problem from comments in the market. Gathering information about a bug in such a situation is a manual process: checking logs, checking versions, manually reproducing the bug, reconfiguring the environment. If there is no clear habit of keeping a list of release changes, logging what could have broken, or at least performing regular smoke tests, the process turns into a repetition of the same steps: check on an emulator (or simulator), run on an old device, update the SDK, clear the cache.

In the support phase, the routine does not disappear – it changes form. Tracking crashes, analyzing feedback, updating dependencies, preparing hotfixes – all of this takes time. Often, the team does not have a clear structure for processing incoming information: someone reads reviews in the market, someone looks at Crashlytics, and someone learns about a bug from a colleague on Slack. Systematizing this work is considered an indirect task, and there are not enough resources for it.

All these problems have a common source: decisions to reduce routine work are made on a case-by-case basis rather than as part of a systematic vision. Even if a single task is automated or simplified, this does not guarantee that the next one will be solved in the same way. Without a coordinated approach, an update process, documentation, and internal accountability, solutions remain vulnerable. And even the best ones are forgotten over time, giving way to the next iteration of manual work, which is faster because it is familiar.

4. Impact of repetitive tasks on development efficiency and quality

«Routine» in development is not always perceived as a problem – it is often disguised as familiar actions that simply need to be done. But this is precisely where its insidiousness lies: these actions do not cause direct resistance, are not considered mistakes, are not recorded in Jira as bugs, and yet they consume time, energy, and focus every day.

As a result, not only speed suffers, but also the quality of development and, just as important, the state of the team itself.

First of all, routine directly affects productivity. When a developer spends an hour every day repeating the same actions – building, updating dependencies, copying templates, manual testing, checking configurations – that hour is lost for tasks that have strategic or creative value. These may seem like minor issues,

but over the course of a week, they add up to half a day, and over the course of a month, several full working days. In projects without well-coordinated automation, such losses become a background load that everyone ignores – but which constantly slows down the overall pace of work.

These repetitive actions also reduce product quality. When a team is working under deadline pressure, any complication, such as manually creating a build or checking settings, forces them to sacrifice other things: test coverage, refactoring, or even basic change verification. If the build fails the first time, the release is often postponed or released with temporary fixes that are then forgotten to be removed. This approach creates technical debt that does not accumulate dramatically but eats away at the code base from within.

At the team level, routine leads to uneven workloads. Often, there are one or two people who know how to assemble a release, configure CI, or quickly check staging. Everyone else waits until they are free from other tasks. The result is dependence on specific people, bottlenecks in the process, and a loss of flexibility. If such an engineer takes a vacation or leaves, the team loses a whole area of knowledge that has to be rebuilt from scratch.

Another important consequence is burnout. Constantly performing monotonous tasks that do not involve any challenge or intellectual effort reduces engagement. A developer who, instead of solving an interesting problem, spends an hour setting up the environment or manually compiling a build, quickly loses motivation. This is especially noticeable in startups or small teams where the number of people is limited, and the tasks are varied. Where there is no system, routine falls on the shoulders of the most active members, and at some point, they simply burn out.

Equally important, routine creates a false impression of efficiency. When a process is built on dozens of manual actions, the team feels that everything is under control because everyone knows what to do and when to do it. But in reality, control becomes illusory: all it takes is a new app store policy or a change in tools, and the system falls apart because nothing is documented, nothing is standardized, and every action is only in the head of a specific team member.

Ultimately, even with good intentions, the lack of a systematic approach means that any implementation – whether CI, code templates, or a build tool – remains fragmented. Without a shared vision and support, such initiatives disappear along with those who initiated them. And even if the next team decides to do things right, it will repeat the same cycle: manual work, fragmented improvements, oblivion, manual work.

Conclusions

An analysis of the daily practices of mobile developers has shown that routine tasks not only remain present in modern workflows, but also have a significant impact on development speed, product quality, and team collaboration. Despite the availability of technical solutions such as automation, templates, and CI/CD tools, most attempts to reduce repetitive tasks remain isolated. They arise as a response to local difficulties, are not integrated into the overall context of the development lifecycle and are not supported by long-term practices.

The problem is further complicated by the fact that much of the routine is perceived as something normal, something that is supposed to be that way. In reality, however, this is one of the hidden barriers to scaling, effective onboarding of new team members, and ultimately product stability. The lack of common standards, transparent processes, and responsibility for maintaining even the simplest automations leads to their gradual decline and, consequently, a return to manual execution.

This study did not aim to provide a comprehensive solution, but rather to identify the problem area, describe it, and structure it. The result was the formation of a comprehensive vision of which stages of mobile development are accompanied by repetitive tasks, how these tasks manifest themselves in a real environment, why the implemented solutions often do not work, and what consequences this has for the team.

Further research should focus on building a practical framework that would allow typical routine tasks to be matched with specific tools, approaches, or organizational practices that have proven effective in a similar context. Such a framework should not only catalog existing solutions but also create a convenient structure for their evaluation, adaptation, and scaling. The focus should be not only on automation, but also on the viability of solutions: how easy they are to implement, maintain, transfer within a team, and apply in different conditions. Over time, such a model could become the basis for unifying approaches to routine tasks in mobile development, which would significantly increase the predictability, stability, and productivity of teamwork.

References

1. Build your app from the command line. Android Developers. URL: <https://developer.android.com/build>

2. About pull requests. GitHub Docs. URL: <https://docs.github.com/en/pull-requests/collaborating-with-pull-requests/proposing-changes-to-your-work-with-pull-requests/about-pull-requests>
3. Add build dependencies. Android Developers. URL: <https://developer.android.com/build/dependencies>
4. Fundamentals of testing. Android Developers. URL: <https://developer.android.com/training/testing/fundamentals>
5. What is Software Automation? GeeksForGeeks. URL: <https://www.geeksforgeeks.org/what-is-software-automation>
6. Create app templates. Android Developers. URL: <https://developer.android.com/studio/projects/templates>
7. Santos I., Filho J., Souza S. A survey on the practices of mobile application testing. Proceedings of the XLVI Latin American Computing Conference. 2020. P. 232–241. URL: <https://doi.org/10.1109/CLEI52000.2020.00034>
8. Patidar A., Suman U. A Survey on Mobile App Development Approaches with the Industry Perspective. International Journal of Open Source Software and Processes. 2022. Vol. 13(1). P. 1–17. URL: <https://doi.org/10.4018/IJOSSP.300754>
9. Berihun N., Dongmo C., Van der Poll J. The Applicability of Automated Testing Frameworks for Mobile Application Testing: A Systematic Literature Review. Computers. 2023. Vol. 12(5). P. 97. URL: <https://doi.org/10.3390/computers12050097>
10. Linares-Vásquez M., Bernal-Cárdenas C., Moran K., Poshyvanyk D. How do Developers Test Android Applications?. Proceedings of the 2017 IEEE International Conference on Software Maintenance and Evolution (ICSME). 2017. P. 613–622. URL: <https://doi.org/10.1109/ICSME.2017.47>

COERCIVE MEASURES APPLIED UNDER THE LEGAL EMERGENCY REGIME

Vypriyskyi A.

The Ukrainian statehood is compelled to assert its right to self-defense in response to the unprovoked armed attack by the Russian Federation utilizing the territory of the Republic of Belarus. The aggressive military tactics employed by the aggressor nation necessitate Ukraine's comprehensive utilization of all available means of resistance [1, c. 153].

Since the onset of the full-scale invasion, the Ukrainian government and allied states have confronted significant challenges, addressed in part through the imposition of martial law nationwide. Warfare encompasses not only frontline combat but also countering the intelligence activities of hostile groups, fighting by collaborators, maintaining public safety and law and order, security and traffic regulation in the rear, preventing and detecting cyberattacks, organizing and carrying out the evacuation of the civilian population from danger, social support of displaced persons, etc. The second year of the war revealed certain problems, the resolution of which posed certain difficulties, thereby necessitating legislative changes, particularly in the areas of criminal, administrative, disciplinary, and civil liability, as well as the adoption and formal establishment of new compulsory legal enforcement procedures regarding the implementation of restrictive measures of the legal regime of martial law. Some of these provisions are completely new, provoked by the extreme conditions of martial law, therefore need in the shortest time a thorough scientific analysis and substantiation in order to predict further consequences and possible improvement, to obtain clear mechanisms of implementation and application in practice.

The aim of this work is to identify the characteristics and scope of administrative coercion during a state of emergency, and to formulate evidence-based suggestions for enhancing both the legislation and its practical application in this domain.

The research relies on a collection of scientific methods and techniques to achieve its objectives, driven by a systemic approach aimed at investigating problems through the integration of their social context and legal framework. The primary method employed is the general scientific dialectical method, facilitating the analysis of current legislation regarding the application of administrative coercion measures during a state of emergency. Additionally, specific methods of scientific inquiry are employed. For instance, the comparative legal method is utilized to explore issues such as the positioning of a state of emergency within the framework of emergency administrative and legal regimes, the responsibilities of various stakeholders in implementing emergency measures, and the normative foundations governing the use of coercion during a state of emergency. Furthermore, systemic and functional methods are employed to analyze the structure of entities responsible for implementing state of emergency measures.

Today's reality is marked by significant changes in the life of the country, accompanied by a range of man-made disasters, natural and man-made emergencies, and socio-political phenomena including mass events and attempts to seize state power, which often lead to group violations of public order and mass riots. The progression of these events directly influences the day-to-day functions of state authorities and local self-government bodies, shaping their operational

peculiarities and modes of interaction. Amidst these challenges, there arises an urgent need for the socialization of all spheres of public life to ensure reliable public order during emergencies of varying natures and scales, prompting the declaration of special states of emergency.

The grounds and procedures for implementing states of emergency and martial law, along with the regulations governing these states, establish a specific framework for the functioning of executive and local self-government bodies. In addressing the consequences of such states, and working to eliminate their causes and conditions, these bodies undertake a series of administrative and legal measures. Foremost among these measures is administrative coercion, wielded exclusively by state authorities and local self-government bodies.

The exploration of theoretical and practical challenges related to the implementation of administrative coercion measures during a state of emergency represents a central and prioritized focus of contemporary scientific inquiry within the realm of state authorities and local self-government. This area of study demands a comprehensive integration of knowledge across various branches of legal regulation.

It is noteworthy that following Ukraine's independence, discussions regarding the utilization of administrative coercion measures during states of emergency did not receive adequate attention within legal scholarship. Instead, most theoretical and practical developments centered around themes such as economic stability, the application of administrative responsibility, the maintenance of public order, and the organization of internal affairs activities under such conditions.

The originality of the research was driven by the imperative to address several key objectives: establishing the position of the legal regime of a state of emergency within the framework of emergency administrative and legal regimes; systematizing the legal provisions governing the use of administrative coercion during a state of emergency; describing the structure of entities responsible for implementing administrative coercion measures during a state of emergency and elucidating the specifics of their authority during this period; classifying various types of administrative coercion applicable during a state of emergency and providing a comprehensive understanding of their essence; clarifying the nuances associated with implementing both primary and supplementary administrative coercion measures in the context of a state of emergency; proposing avenues for enhancing the legal framework governing the application of administrative coercion measures during a state of emergency.

These objectives collectively contributed to the originality and significance of the research, addressing crucial aspects of administrative coercion within the unique context of a state of emergency.

Introduction

Legal regulation of relations that arise during natural disasters, man-made accidents and catastrophes, socio-political and military conflicts, has historically been carried out using state acts that declare the issue of applying legal measures of an extraordinary nature [1, p. 56].

The practical application of administrative coercion measures is always associated with the restriction, often quite significant, of the constitutional rights and freedoms of citizens, as well as the rights and legitimate interests of legal entities. In this regard, it is very relevant and expedient to create effective legal foundations for the application. The effectiveness of the aforementioned activities in many cases

also depends on the perfection of the regulatory regulation of the process of activity of state and law enforcement bodies regarding the application of administrative coercion measures, the clarity of regulations, and the presence of a developed system of legislation. These measures, that is, to ensure their proper legal implementation. In the theory of state and law, legal regulation is considered to be an effective, normative and organizational influence exerted on social relations through a system of legal means (individual regulations, norms, legal relations) with the aim of their protection and development in accordance with social needs [2, p. 89].

The relevance of this work lies in the fact that, on the one hand, in the science of administrative law there is great interest in administrative coercive measures, on the other hand, laws relating to the legal regulation of coercive measures in conditions of emergency and martial law, due to their limited application, are practically imperfect.

The purpose of this work is to study the problems of theoretical improvement, legal regulation and practical application of administrative coercion measures in a state of emergency. Taking into account the set goal, the following tasks should be solved:

- analysis of research on the application of administrative and legal measures in a state of emergency;
- to determine the concept, features and certain content of the legal institution of social relations in the field of application of administrative and legal measures in a state of emergency;
- to develop recommendations aimed at improving the theoretical and legal foundations of administrative coercion under the conditions of legal regimes of emergency status.

Analyzing the state of research into issues related to the legal nature of administrative and legal measures in a state of emergency, at different times these problems were addressed by A.T. Komzyuka, V.K. Kolpakova, T.O. Kolomoyets, A.V. Basov, O.K. Bezsmertny, Yu.V. Drozd, O.K. Bezsmertny, M.I. Yeropkina, R.V. Kisil, B.V. Rossinsky, Yu.M. Starylov, who have prepared scientific works of a monographic nature on this topic. At the same time, the problems associated with the use of coercive measures in the specified circumstances have almost not found their scientific justification in the legal literature. Therefore, with the presence of a significant number of various scientific sources and sources from the mass media, to this there were no scientific works that would explore the full range of extraordinary administrative and legal measures when maintaining the legal regime of a state of emergency.

The content of coercion under the legal regime of a state of emergency

Constitutional and civil rights and freedoms of man and citizen and their guarantees determine the content and direction of the functioning of the state. To fulfill tasks in the field of ensuring the security of the state and citizens from various threats, and in particular emergency situations that lead to the introduction of a state of emergency, the state uses, among other things, various administrative and legal means.

Administrative and legal means establish a mechanism for the implementation of the rights and freedoms of citizens, their implementation is ensured through administrative law enforcement, law enforcement administrative and legal means are used to protect the rights and freedoms of citizens from various kinds of threats (consequences of emergencies, administrative offenses, etc.), the basis of which are measures of administrative coercion.

It should be noted that one of the most dangerous threats to the state and its citizens are emergencies, as their occurrence causes unforeseen material costs, economic burden, human casualties, destruction of infrastructure, and, consequently, an increase in forced unemployment, economic impoverishment, which objectively contributes to the growth of social tension in the country among the population. Regarding the statistics of emergency situations, in 2019, compared to 2018, the total number of emergency situations decreased by twelve percent, while the number of man-made and natural emergencies decreased, while the number of socio-political emergencies increased slightly. Compared to 2018, the number of state-level emergencies remained unchanged, while the number of regional, local, and facility-level emergencies decreased. It should be noted that the number of regional emergencies decreased by two and a half times. During 2019, seven state-level emergencies were registered (in 2018, three state-level man-made emergencies).

In addition to natural and man-made emergencies, the Ukrainian people are prone to the emergence of rebellious socio-political emergencies. This is evidenced by the events that took place in 1998 in Kyiv, during the opposition's actions against the existing government, and in 2004–2005, the organization and conduct of the so-called «Orange Revolution», as well as political and social changes in Ukraine from November 21, 2013 to February 2014, which were caused by resistance to the departure of the country's political leadership from the legally established course of European integration and subsequent opposition to this course. The armed conflict in the region began in mid-April 2014, when armed groups of pro-Russian activists seized administrative buildings and police stations in Sloviansk, Artemivsk, and

Kramatorsk. During the events described above, state authorities used certain legal, organizational, financial measures that should urgently ensure stabilization of the situation. However, the practice of applying these measures shows that the current legislation in regulating many issues lags behind the urgent needs of practical application, it does not meet the needs of today, has a number of shortcomings, legal gaps and conflicts, does not define a clear concept of certain legal categories, which complicates the law enforcement activities of relevant state bodies and local self-government bodies.

This position is supported by the results of a survey of civil servants, among whom more than seventy-six percent are convinced of the need to change the regulatory framework that regulates relations in emergency situations and states of emergency.

The legal basis for the application of administrative coercion measures under the state of emergency regime is the norms enshrined in:

- the Constitution of Ukraine (Article 64, Clause 31 of Article 85, Clause 19 of Article 92, Clauses 20, 21 of Article 106, Clause 10 of Article 138);
- the Civil Protection Code of Ukraine;
- Laws of Ukraine «On the Legal Regime of the State of Emergency», «On the National Security and Defense Council», «On the National Police», «On the Internal Troops of the Ministry of Internal Affairs of Ukraine», «On Mobilization Training and Mobilization», «On the Security Service of Ukraine», «On the Armed Forces of Ukraine»;
- relevant Decrees of the President;
- Resolutions of the Cabinet of Ministers of Ukraine regulating issues arising in the specified conditions;
- orders, instructions, rules issued by relevant state and law enforcement agencies.

When a state of emergency is introduced in a certain territory, emergency measures acquire special importance, which are a combination of administrative and legal prohibitions and obligations for individuals and legal entities and the granting of additional rights to state authorities. These measures are a necessary component of the system of elements of a state of emergency, and their legal form is enshrined in Articles 16, 17, 18, 21 of the Law of Ukraine «On the Legal Regime of a State of Emergency» [3].

The introduction of these measures entails significant changes in the legal status of individuals and legal entities, and therefore the transfer to executive authorities and local self-government bodies of extraordinary powers, which they

are granted to eliminate the security threat that has arisen in a certain territory. At the same time, the competence of the state apparatus is expanded, mainly by restricting the rights and freedoms of citizens and legal entities. In the event of the introduction of a state of emergency, a special enhanced regime of activity of state authorities and administration arises, which allows restricting some constitutional and civil rights and freedoms of citizens and the rights and interests of legal entities.

Measures applied under the conditions of a state of emergency are unacceptable and unconstitutional in normal life conditions, but their application is necessary in the event of an emergency situation that necessitates the introduction of a legal state of emergency [4, p. 71].

The application by the authorities of restrictions on the rights of individuals and legal entities is a forced measure, because to eliminate the security threat to both the state and citizens, unusual legal means are used in the normal conditions of life in the country.

The negative consequences caused by the emergency situation require executive authorities and local self-government bodies to apply adequate measures specified in the Law of Ukraine «On the Legal Regime of the State of Emergency» otherwise, the state is unable to ensure the fulfillment of the functions assigned to it, which are enshrined in the Basic Law of the state, namely, ensuring security in society.

An analysis of the Law of Ukraine «On the Legal Regime of a State of Emergency» shows that a significant place in its content is given to determining the procedure for applying administrative coercion measures under the conditions of a state of emergency [3]. Note that one of the mandatory conditions for introducing a state of emergency on the territory of Ukraine or a separate part of it is the definition in the Decree of the President of Ukraine «On the Introduction of a State of Emergency» a list and limits of emergency measures, an exhaustive list of constitutional rights and freedoms of man and citizen that are temporarily restricted in connection with the introduction of a state of emergency, as well as a list of temporary restrictions on the rights and legitimate interests of legal entities, indicating the period of validity of these restrictions.

The state of emergency significantly expands the scope of legal regulation of individuals regarding the exercise of their rights and freedoms. The administrative apparatus is given the authority to restrict or prohibit the exercise of certain constitutional rights that are protected under normal conditions.

By adapting the types of legal regulation, the most effective opportunity to control the situation is achieved. Citizens should act as necessary at a specific

moment of crisis – everything that is permitted by law is permitted. Thus, the state of emergency creates a certain direction in the realization of subjects' rights and freedoms. Our study, conducted in the course of our research, studying the practical experience of employees of the State Emergency Service during the application of coercive measures, revealed that the absolute majority of rescuers face resistance from the population.

The considered interrelations of the types of legal regulation create a special «restrictive mood» in the mechanism of legal influence. Administrative and legal prohibitions are put forward, which allow directing the behavior of citizens into a «rigid administrative and legal channel». We are faced with the use of administrative coercion, which is regulated by administrative law, where it has been most fully investigated and its regulatory basis has been determined [5, p. 153].

Clarifying the issue of the procedure for applying administrative coercion by state bodies under the conditions of a state of emergency involves, first of all, an analysis of scientific and practical views on the definition of the concept and features of administrative coercion.

The works of administrative scientists (A.T. Komzyuk, V.K. Kolpakova, T.O. Kolomoyets, etc.) have examined the concept of administrative coercion in detail, in this regard, we will focus on a general analysis of this definition.

Some lawyers consider administrative coercion as a system of means of psychological or physical influence on the consciousness and behavior of people in order to achieve clear fulfillment of established duties, develop social relations, and ensure law and order [6, p. 204].

Ukrainian researchers, such as: V. K. Kolpakov, O. V. Kuzmenko, note that administrative coercion is a power-based, unilaterally implemented and applied in cases provided for by legal norms, on behalf of the state to the subjects of legal relations, firstly, preventive measures to prevent offenses, secondly, measures to stop offenses, thirdly, measures of responsibility for violation of prohibitions [7, p. 193].

T.O. Kolomoyets in her dissertation research «Administrative coercion in the public law of Ukraine; theory, experience and practice of implementation» proposed her own definition of administrative coercion in the public law of Ukraine, which should be understood as a special, complex, a polystructural type of state-legal coercion, i.e. methods of official physical or psychological influence of government bodies on individuals and legal entities, defined by public law norms, in the form of personal, property, organizational restrictions on their rights, freedoms and interests in cases of commission of illegal acts by these persons (in the field of relations of a public nature) or in extraordinary circumstances within the framework

of a separate, economical, simplified, operational administrative proceeding in order to achieve a multi-faceted retro-perspective goal (prevention, cessation of illegal acts, prevention and localization of the consequences of emergency situations, ensuring the proceedings in cases of offenses, bringing guilty persons to justice) [9, p. 63].

The most successful definition of administrative coercion is given by the author A.T. Komzyuk in the monograph «Measures of administrative coercion in law enforcement activities of the police: concept, types and organizational and legal issues of implementation». The scientist notes that administrative coercion is the application by relevant subjects to persons who are not under their control, regardless of the will and desire of the latter, of measures of influence of a moral, property and other nature provided for by administrative and legal norms in order to protect social relations arising in in the field of public administration, through the prevention and cessation of offenses, punishment for their commission [10, p. 214].

It is worth noting that this definition is also characteristic of administrative coercion, which is used in conditions of emergency and martial law, but it has its own characteristics. Analysis of scientific, educational, and journalistic sources allows us to trace a certain trend - administrativeists formulate definitions of administrative coercion, and then characterize the most important, most significant features. The definition alone, together with the study of its features, provides a general idea of administrative coercion, and in particular, of administrative coercion applied under the legal regimes of emergency and martial law.

It should be noted that administrative coercion applied under the conditions of the state of emergency regime is characterized by all the features of administrative coercion, namely: official, state-authority character; multiplicity and diversity of subjects of application (state bodies, and in some cases, public organizations); number of persons to whom the application is made; lack of official subordination between the subjects of application and the persons to whom the relevant application is made; specifics of the legal and factual grounds(both the presence and absence of an illegal act); multivariate external forms of manifestation; coercive nature; simplicity, efficiency, economy of the procedural regime of application; administrative coercive measures are established, changed, canceled by administrative acts depending on the need [9, p. 211].

It should be noted that administrative coercion, which is applied under the conditions of a state of emergency, is characterized by certain features. Firstly, administrative coercion measures under the state of emergency are applied only for the period specified in the Law of Ukraine «On the Legal Regime of the State of Emergency», namely up to thirty days on the territory of Ukraine, and up to

sixty days in some of its regions. That is, administrative coercion applied under the state of emergency is temporary in nature.

Secondly, executive authorities and local self-government bodies, under the conditions of the state of emergency, in order to eliminate the consequences of the emergency situation, may apply only those measures of administrative coercion that are directly provided for in the Presidential Decree of Ukraine «On the introduction of a state of emergency» and in the Law of Ukraine «On the legal regime of a state of emergency». It should be noted that during the elimination of an emergency situation, the occurrence of which led to the introduction of a state of emergency, state bodies may apply other measures of administrative coercion, however, their use is typical for the daily activities of these bodies. Therefore, the procedure for their use is not regulated by the Law of Ukraine «On the Legal Regime of a State of Emergency». Thus, in the event of mass unrest or other emergency situations of a socio-political nature, employees of internal affairs bodies may use physical force, special means, and weapons in order to resolve this situation. However, the legal grounds for applying these measures are directly provided for in Articles 42, 43, 44, 45, 46 of the Law of Ukraine «On the National Police».

Thirdly, the effect of administrative coercion measures applied under the conditions of the state of emergency applies to individuals and legal entities who are located or reside in the territory of the special regime.

The above analysis of the concept and features of administrative coercion makes it possible to use the concept of administrative coercion, which is applied under the conditions of the state of emergency, in particular, its application by relevant entities to individuals and legal entities that are not under their subordination, on the territory and during the state of emergency, temporary, special restrictive measures of moral, economic and other nature are imposed in order to protect national security.

This definition of administrative coercion, which is applied in the conditions of the state of emergency, has several positive aspects. First, the objective necessity of administrative coercion in the conditions of the state of emergency to solve the tasks of ensuring the activities of some state bodies (for example, internal affairs bodies) is clear, ensuring the security of citizens, the normal functioning of the national economy, state authorities and local self-government bodies, and the protection of constitutional order. This is the essence of administrative coercion in the process of interaction between state authorities and individuals and legal entities. Secondly, with this understanding of administrative coercion applied

under the conditions of the state of emergency, its content is relatively easy to separate from other administrative and legal means of coercion in the aspect of legal regulation. Thirdly, the presented understanding of administrative coercion applied under the conditions of the state of emergency allows us to fundamentally re-examine the application of coercive measures in the mechanism of the state of emergency. Fourth, indicating the understanding of administrative coercion applied under the conditions of the state of emergency regime makes it possible to understand the nature of relations and the specifics of this subject of regulation.

The content of coercion under the conditions of a state of emergency

When analyzing the content of administrative coercion applied in the aspect of the state of emergency, it is necessary to examine a very debatable issue, namely the grounds for the application and classification of administrative coercion measures.

Today, there are several views on the grounds for applying administrative coercion measures. One group of scientists believes that administrative coercion is applied only to persons who have committed administrative offenses, The second group of scientists believes that administrative coercion can be used not only as a reaction of the state to offenses, but also, if necessary, to protect public order and ensure public safety in special and extraordinary conditions. Let's analyze the above.

Representatives of the first group believe that administrative coercion necessarily has the character of administrative punishment and is carried out only in connection with an unlawful and socially dangerous act. R.V. Kisil, B.V. Rossinsky, Yu.M. Starylov note that state coercion is carried out as a reaction of state bodies to unlawful, socially dangerous behavior of people. Its application is due to the conflict between the state will, which is expressed in a legal act, and the individual will of the persons who violate it.

The view of the above-mentioned scholars regarding the restriction of citizens' rights and freedoms in the context of an emergency situation is debatable. These measures are not coercive measures, regardless of the «unfavorable» consequences that occur for the person. Not all «unfavorable» consequences, even if they are provided for by law and arise as a result of conscious actions of individuals, are the result of coercion. Coercion applies only to a specific legal subject who has violated the rules of law and has a personified individual character. The same measure may be coercion in one case and not in another, depending on whether the behavior was unlawful. Therefore, the use of administrative coercion is associated with the presence of an administrative offense [11, p. 32].

However, we consider the positions of those scientists who insist that administrative coercion measures can be applied to both offenders and persons who have not committed an offense to be convincing [12, p. 31]. As A.V. Basov aptly notes, the basis for lawful restriction of the administrative and legal status of citizens, in addition to offenses, can also be legally significant events, such as emergencies of a technogenic or natural nature, social conflicts, military actions and armed riots, etc. [5, p. 162].

Scientist O.K. Bezsmertny in his dissertation research «Administrative preventive measures applied by internal affairs bodies» notes that administrative coercion as a whole performs the tasks of protecting, developing and strengthening normal administrative relations and eliminating offenses, eliminating their consequences that may harm public or state interests. It is a legal means of protecting public relations from unlawful actions, restoring the lawful state, ensuring the protection of public order and ensuring public safety in normal conditions of life and in the event of emergency situations [13, p. 175].

A rather interesting definition of administrative coercion was provided by Yu.V. Drozd. In his opinion, the criterion of administrative coercion is the method of administrative influence of the state in the person of state bodies authorized by law, officials on subjects of social life (individuals and legal entities) in order to prevent and stop illegal behavior in the interests of protecting the rights and legitimate interests of all members of society, ensuring public safety and law and order, and bringing those responsible to legal responsibility [14, p. 94].

As S.O. Kuznichenko aptly notes, among the factors requiring the use of coercive measures in situations not related to the commission of offenses, there may be a state necessity or a social necessity to prevent the occurrence of harmful consequences, protect public order and ensure security. Under such conditions, authorities, using the additional powers granted to them, restrict the possibility of citizens to exercise certain constitutional rights and freedoms. At the same time, this restriction (coercion) is not related to the commission of illegal actions, but is applied in connection with the occurrence of a relevant event [15, p. 79].

Administrative coercion during the state of emergency is the basis of legal regulation of social relations, as it is immersed in various social relations: civil, constitutional, economic, environmental, labor, administrative, etc. In order to eliminate the impact of the negative factor on the population in the relevant territory as soon as possible, emergency legal measures are applied, which cause emergency restrictions on the rights and freedoms of individuals and business entities.

The lawful restriction of the administrative and legal status of citizens, acting as the implementation of the state of emergency regime, became the basis for the allocation of a special form of administrative and legal regulation. Its specificity lies in the nature of legal influence in the specified conditions and is manifested in the establishment of a regime of administrative and legal restrictions. The expansion of the competence of administrative bodies makes it possible to change the nature of coercive influence to a certain extent. The effect of regulations on the restriction of rights and freedoms under the conditions of the state of emergency does not apply to specific individuals, but to the population of the entire territory where the regime operates. In emergency conditions, this type of regulation becomes convincingly predominant [5, p. 169].

An idea of administrative coercion is impossible without analyzing its classification. Recently, this issue has been the focus of great attention of scientists, as evidenced by a number of studies. However, the issue of classifying administrative coercive measures under the state of emergency has not been finally defined in relevant scientific works, despite the significant practical and theoretical importance of clarifying this issue. As O. K. Bezsmertny aptly notes, a clear classification of measures is necessary, first of all, to clarify the essence of various coercive measures applied by executive authorities and local self-government bodies under the conditions of the state of emergency and martial law, to understand the purpose, «legal potential», their correlation and interaction [13, p. 215].

Some scientists believe that all administrative coercion measures should be divided into: liability measures and protective measures. The first group includes enforcement measures, the second – measures of a preventive nature [13, p.152].

Other lawyers express the opinion that administrative coercion does not include administrative preventive measures, arguing that administrative preventive measures are not coercive measures, since they belong to prohibitive norms of law. These norms are not of a personal nature and are applied to an individual, but are addressed to all citizens without exception, therefore they cannot be considered as a measure of administrative coercion. After all, only an individual act of management that has a specific addressee can be such [17, p. 35]

In our opinion, the position of M.I. Yeropkin is very correct, who believes that administrative coercion measures should be divided into the following measures of influence: administrative termination measures; administrative preventive measures, administrative penalties. The proposed classification is supported and used in their works by most modern administrative scientists: V.B. Averyanov, O.M. Bandurka, O.K. Bezsmertny, V.K. Kolpakov, A.T. Komzyuk, T.O. Kolomiets, etc. They believe

that the above groups of coercive measures are applied in cases of: bringing to administrative responsibility persons who have committed administrative offenses, stopping illegal actions, preventing offenses or ensuring security as a result of a natural disaster or other violations of any spheres of society's life.

Regarding the classification of administrative coercive measures provided for by the Law of Ukraine «On the Legal Regime of a State of Emergency» and which are applied under the conditions of this regime, we note that, according to A.V. Basov, some emergency situations caused by the introduction of a state of emergency are unlawful (terrorist acts, mass riots, crossing the state border, etc.) and their elimination requires the country's authorities to take both administrative preventive measures and administrative termination measures.

Taking this into account, A.V. Basov believes that the measures provided for by the Law of Ukraine «On the Legal Regime of a State of Emergency» relate to both administrative preventive measures and administrative termination measures. In the event of the introduction of a state of emergency in a certain territory, the Presidential Decree «On the Introduction of a State of Emergency» will necessarily specify additional measures that will be applied during the state of emergency and are aimed both at preventing the occurrence of unlawful behavior and at its cessation [5].

Of course, this view has the right to exist, but the position of T.O. Kolomoiets is well-founded and appropriate, who believes that depending on the factual grounds for introducing a state of emergency, all measures applied by state authorities are aimed primarily at protecting people: preventing death and mutilation [9, p. 341].

Administrative warning measures applied under the state of emergency can be conditionally divided into three groups, combining those with similar legal and procedural properties.

The first group is general administrative and coercive measures of a preventive nature of the legal regime of a state of emergency: strengthening the protection of civil order and facilities that ensure the vital activity of the population; banning strikes; restrictions on the movement of vehicles and their inspection; establishment of a special regime of entry and exit, as well as restrictions on freedom of movement in the territory where a state of emergency is introduced; prohibition of holding mass events, except for events the prohibition of which is established by the court.

The second group is preventive administrative enforcement measures in a state of emergency due to man-made or natural emergencies: evacuation of people from places that are dangerous to live in; introduction of a special procedure for the distribution of food and essential items; removal from work for the period of the state of emergency, in case of improper performance of their duties, heads of state

enterprises, institutions and organizations on whose activities the normalization of the situation in the area of the state of emergency depends; establishment of housing obligations for legal entities for temporary accommodation of evacuated or temporarily resettled population; establishing quarantine and carrying out other mandatory sanitary and anti-epidemic measures.

The third group is administrative coercive measures of a preventive nature under the legal regime of a state of emergency in connection with mass violations of civil order: the introduction of a curfew (a ban on being on the streets and in other public places without specially issued passes and identity cards at established hours of the day); checking citizens' documents, and if necessary, conducting a personal search, inspection of belongings, vehicles, luggage and cargo, office premises and housing of citizens; restrictions or temporary bans on the sale of weapons, poisonous and potent chemicals, as well as alcoholic beverages; special rules for the use of communications and the transmission of information via computer networks; temporary confiscation of registered firearms and cold weapons and ammunition from citizens, etc.

All other measures applied by state bodies under the conditions of the state of emergency and not provided for by the Law of Ukraine «On the Legal Regime of Emergency» can be attributed to other types of administrative coercion measures (administrative detention, use of physical force, use of firearms, etc.), however, most of them are used in everyday life and ordinary living conditions.

It is significant that all administrative coercive measures applied under the state of emergency can be classified according to certain criteria developed in administrative and constitutional law. Thus, depending on the purpose of application, these measures can be divided into: measures of general preventive nature (establishing a special regime of entry and exit, as well as restrictions on freedom of movement in the territory where a state of emergency is introduced; restrictions on the movement of transport and its inspection); measures used to prevent administrative offenses by specific individuals (strengthening the protection of public order and facilities that ensure the vital activity of the population; banning mass gatherings, except for events the ban on which is established by the court; banning strikes).

Depending on the nature of the law enforcement impact, administrative coercive measures applied during a state of emergency are divided into: personal (introduction of a special procedure for the distribution of food and essential items), property (establishment of housing obligations for legal entities for temporary accommodation of evacuated or temporarily resettled population, emergency and

rescue units and military units involved in overcoming emergency situations) and organizational and legal (temporary ban on the construction of new, expansion of existing enterprises and other facilities whose activities are not related to the elimination of an emergency situation or ensuring the vital activities of the population and emergency rescue units). Depending on the object of influence – those applied to individuals (temporary or irreversible evacuation of people from places that are dangerous for living, with the mandatory provision of them with permanent or temporary housing; prohibition of conscripts and those liable for military service from changing their place of residence without the knowledge of the relevant military commissariat); curfew; checking citizens' documents, and if necessary – conducting a personal inspection, inspection of things, cars, luggage and cargo, office premises and housing of citizens; legal entities (changing the operating mode of enterprises, institutions, organizations of all forms of ownership, reorienting them to the production of products necessary in the state of emergency, fulfillment of mobilization tasks and government orders, other changes in production activities necessary for carrying out emergency rescue and restoration work; removal from work for the period of the state of emergency, in case of improper performance of their duties, of heads of state enterprises, institutions and organizations on whose activities the normalization of the situation in the area of the state of emergency depends, and the assignment of temporary performance of the duties of the said heads to other persons); mixed (mobilization and use of resources of enterprises, institutions and organizations, regardless of the form of ownership, to avert danger and eliminate emergencies with mandatory compensation for losses incurred; restrictions or temporary bans on the sale of weapons, poisonous and potent chemicals, as well as alcohol and alcohol-based substances).

By implementation period – those that are implemented by performing certain one-time actions (checking citizens' documents, and if necessary, conducting a personal inspection, inspection of things, vehicles, luggage and cargo, office premises and housing of citizens), not related to the period (prohibition of the production and distribution of information materials that may destabilize the situation), those that are distinguished by the duration of action (introduction of a curfew).

Conclusions

Administrative coercion applied under the conditions of the state of emergency is the application by the relevant bodies of state power and local self-government (to individuals or legal entities not under their control, on the territory and during the state of emergency and martial law, temporary, special restrictive measures

of influence of a moral, property, personal and other nature in order to protect national security and ensure the safety of life of the population of the country and other countries of the civilized world.

This definition of administrative coercion applied under the conditions of the state of emergency has several positive aspects. First, with this understanding of administrative coercion applied under the conditions of the state of emergency, its content is relatively easy to distinguish from other administrative and legal means of coercion in the mechanism of legal regulation. Secondly, the objective need for administrative coercion under the conditions of a state of emergency and martial law is clear, to solve the tasks of ensuring the activities of certain state bodies (for example, the Security Service of Ukraine, the Armed Forces of Ukraine, the State Border Service of Ukraine, the Ministry of Internal Affairs), ensuring the security of individuals, the normal functioning of the national economy and its development, cybersecurity, energy security, state authorities and local self-government bodies, and protecting the constitutional order. This is the advantage of administrative coercion in the process of interaction of state bodies with citizens and legal entities. Thirdly, indicating the understanding of administrative coercion, which is applied in the conditions of the state of emergency regime, to understand the nature of relations, the specifics of the subject of regulation. Fourth, the presented understanding of administrative coercion applied under the conditions of the state of emergency and martial law allows us to show for the first time in a fundamentally new way the use of administrative coercive measures in the mechanism of the state of emergency.

References

1. Kuznichenko S.O. (2001). Management of internal affairs bodies in special conditions caused by anomalous man-made and natural phenomena. – *H.: Nation Edition. University of Internal Affairs*, pp. 170. (in Ukrainian)
2. V. M. Marchuk, L. V. Nikolayeva, T. O. (2019) Theory of the state and law: a textbook – Kyiv: Kyiv. national trade and economy University, pp. 424. (in Ukrainian)
3. On the legal regime of the state of emergency: Law of Ukraine (2000) no. 1550-III dated // VVR of Ukraine. no. 23. – art. 176. (in Ukrainian)
4. Basov A.V. Administrative coercion in the context of a state of emergency (2005) Actual problems of modern science in the research of young scientists. – Simferopol: 2005. – no 8. p. 141. (in Ukrainian)
5. A.V. Basov Administrative and legal regime of state of emergency (2007) *Diss.... can. legal Sciences: 12.00.07 – Kh.*, – pp. 201. (in Ukrainian)
6. Battle Yu.P. Administrative law of Ukraine (2005) *Pravo*, pp. 459. (in Ukrainian)

7. Kolpakov V.K., Kuzmenko O.V. Administrative law of Ukraine. (2003) *Yurinkom «Inter Publishing House»* pp. 544. (in Ukrainian)
8. Shestak L.V. Administrative Law. (2011) *Chernihiv State Institute of Law, Social Technologies and Labor* pp. 256. (in Ukrainian)
9. Kolomoets T.O. Administrative coercion in the public law of Ukraine: theory, experience and practice of implementation (2005) dissertation of the doctor of legal sciences: 12.00.07 pp. 455. (in Ukrainian)
10. Komzyuk A.T. Administrative coercion in the law enforcement activities of the police of Ukraine (2002), Doctor of Law dissertation 12.00.07. pp. 407. (in Ukrainian)
11. Kisil Z.R., Kisil R.V. Administrative Law. (2011) *Publishing houses of Alerta; TSU*, pp. 696. (in Ukrainian).
12. Kurochka M.I. Administrative coercion due to the need to stop the offense (2015) *Law Forum*. no. 4. pp. 133–136. (in Ukrainian).
13. Bezsmertnyy E.O. Administrative preventive measures applied by internal affairs bodies (1997) thesis of the candidate of legal sciences 12.00.07 pp. 155. (in Ukrainian).
14. Drozd Yu.V. Correlation of administrative responsibility with other types of administrative coercion (2010) autoref. thesis for the competition of scientists. candidate degree law Sci.: 12.00.07, pp. 18. (in Ukrainian).
15. Kuznychenko S.O. The Law of Ukraine «On the Legal Regime of a State of Emergency»: Science and Practice. comment (2015) *Odessa State University of Internal Affairs* pp. 163. (in Ukrainian).
16. Kuznichenko S. O. Administrative and legal consequences of the introduction of emergency administrative and legal regimes (2008) *Bulletin of the Kharkiv National University of Internal Affairs* no 43 pp. 211–218. (in Ukrainian).
17. Galunko V.V., Olefir V.I., Pikhtin M.P. etc. Administrative law of Ukraine. (2011). *PJSC «Kherson City Printing House»* pp. 320 (in Ukrainian).
18. Averyanov V. B. Administrative Law of Ukraine (2004) *«Yuridichna Dumka» Publishing House*. pp. 584. (in Ukrainian).
19. Kolomoets T.O. Administrative coercion in the public law of Ukraine: theory, experience and implementation practice (2004). *Polygraph*. pp. 404. (in Ukrainian).
20. About the National Police: Law of Ukraine (2015) Information of the Verkhovna Rada, no. 40–41, pp. 379. (in Ukrainian).
21. On the legal regime of martial law: Law of Ukraine (2015) No. 389-VIII. *Voice of Ukraine* no. 101. pp. 311. (in Ukrainian).
22. Decree of the President of Ukraine «On the introduction of martial law in Ukraine». (2022) *Voice of Ukraine* no. 64/2022. pp. 210 (in Ukrainian).
23. Mernyk A.M., Kuzmina V.O., Burlakov V.M. Limitation of human rights and freedoms in modern conditions: theoretical and practical aspects. (2020) *Legal scientific electronic journal*. no. 2. pp. 210 (in Ukrainian).
24. On the principles of state regulatory policy in the field of economic activity: Law of Ukraine (2022) No. 1160. *Voice of Ukraine*. no. 67/2022. pp. 210 (in Ukrainian).

Scientific publication

**INNOVATIVE TECHNOLOGIES
FOR PROJECT AND PROGRAM MANAGEMENT**

*Collective monograph
edited by I. Linde*

**LĒMUMU ATBALSTA SISTĒMAS
PROJEKTU UN PROGRAMMU VADĪBĀ**

*Kolektīvas monogrāfija
I. Linde zinātniskajā redakcijā*

Parakstīts iespiešanai 2025-26-09. Reģ. № 07-10.

Formats 60x84/16 Ofsets. Ofseta papīrs.

14,8 uzsk.izd.l Metiens 300 eks. Pasūt. № 144.

Tipogrāfija «Landmark» SIA, Ūnijas iela 8, k.8, Rīga, LV-1084.

Reģistrācijas apliecības numurs: 40003052610. Dibināts: 28.12.1991.

