

**COMMUNICATION CLASSIFIED AS A WEAPON AND GENERATIVE
ARTIFICIAL INTELLIGENCE AS INFORMATION OVERLOAD:
SITUATION-FORECAST-CONTROL (S-F-C) INDICATORS AND
MULTIMODAL RESILIENCE MATRIX IN THE OXIPO MODEL**

Author(s) / Szerző(k):

Dávid Horváth

National University of Public Service (Hungary)

E-mail:

david@netokracia.hu

**Cite:
Idézés:**

Horváth, Dávid (2025). Communication Classified as a Weapon and Generative Artificial Intelligence as Information Overload: Situation-Forecast-Control (S-F-C) Indicators and Multimodal Resilience Matrix in the OxIPO Model. *OxIPO – Interdiszciplináris tudományos folyóirat*, VII. évfolyam 2025/4. szám. 23-48.

Doi: <https://www.doi.org/10.35405/OXIPO.2025.4.23>



This work is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License.

EP / EE:

Ethics Permission / Etikai engedély: KFS/2025/O4-0008

**Reviewers:
Lektorok:**

Public Reviewers / Nyilvános Lektorok:

1. Éva Jakusné Harnos (Ph.D.), National University of Public Service
2. János Juhász(PhD.), University of Miskolc (Hungary)

Anonymous reviewers / Anonim lektorok:

3. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)
4. Anonymous reviewer (Ph.D.) / Anonim lektor (Ph.D.)

Absztrakt

*FEGYVERMINŐSÉGŰ KOMMUNIKÁCIÓ ÉS GENERATÍV MI, MINT
INFORMÁCIÓFELDOLGOZÁSI TÚLTERHELÉS: HELYZET-ELŐREJELZÉS-
SZABÁLYOZÁS (S-F-C) MUTATÓK ÉS MULTIMODÁLIS RUGALMASSÁGI
MÁTRIX AZ OXIPO MODELLBEN*

A tanulmány tézise, hogy a fegyvernek minősülő kommunikáció (FMK) és a generatív mesterséges intelligencia az emberi információfeldolgozást (lásd: OxIPO/HIP: $O \times (I + P + Out)$) célzottan túlterheli. Bemutatom a saját Helyzet-Előrejelzés-Irányítás (H-E-I; S-F-C) mátrix illesztését, és bevezetem a Multimodális Reziliencia Mátrixot

(MRM) mint oktatási kontroll-eszközt szöveg–kép–hang–videó esetén. Rövid keretes példák (Krim 2014, COVID–19 infodémia, Ukrajna 2022–) validálják.

Kulcsszavak: fegyvernek minősülő kommunikáció (FMK); emberi információfeldolgozás (HIP); Multimodális Reziliencia Mátrix (MRM); generatív mesterséges intelligencia

Diszciplína: pszichológia, hadtudomány, kommunikáció- és médiatudomány

Abstract

The study posits that weapon-grade communication (WGC) and generative artificial intelligence deliberately overload human information processing (OxIPO/HIP: $O \times (I + P + Out)$). I present the integration of the proprietary Situation–Forecast–Control (S–F–C) matrix and introduce the Multimodal Resilience Matrix (MRM) as an educational control tool for text, image, audio, and video modalities. The framework is validated through brief vignettes (Crimea 2014, COVID–19 infodemic, Ukraine 2022–).

Keywords: Weapon-Grade Communication (WGC); Human Information Processing (HIP); Multimodal Resilience Matrix (MRM); Generative AI

Disciplines: Psychology, Military Science, Communication and Media Studies

Introduction: The Cognitive Space as an Operational Area

The relevance of this topic lies in one of the key insights of the 21st-century security environment, namely that the outcome of conflicts is increasingly determined not by physical processes, but by those taking place in the cognitive space. In this space, information is no longer merely a means of providing information, but a tool capable of triggering operational effects: it can be used to distort or overload the target group's perception, interpretation, and decision-making routines. The study describes this phenomenon using the term weapon-grade communication (FMK): it is a deliberately organized, strategic form of influence that cannot be grasped at the level of individual lies, but rather through the

shifting conditions of the information ecosystem (platform architecture, recommendation systems, moderation, media routines, user psychology).

The threat has taken a qualitative leap forward in the past decade: Crimea in 2014 was the prototype of hybrid logic, COVID-19 was the dress rehearsal for global cognitive overload, and the Russian-Ukrainian war that began in 2022 is an example of a scaled, industrial-paced narrative war. Generative artificial intelligence (AI) entered this process by drastically reducing the cost of content production while easily maximizing the appearance of authenticity: the question is not how much misleading content there is, but when quantity turns into quality (thematization, emotional tuning, erosion of public trust).

From the perspective of human information processing (HIP), the study argues that modern FMK/MI does not always seek to persuade, but rather to overload, thereby disrupting the decision-making chain. For this reason, the challenge today is not merely recognition, but the measurable and teachable strengthening of resilience (resistance).

The literature abounds in descriptions and case studies, but frameworks that directly link diagnosis (what is happening?) with intervention (what should we do?) are rarer. This study is therefore methodological: its aim is to construct an operationalizable analytical and defense framework that can later be applied to case studies (Crimea/COVID/Ukraine) in a reproducible manner.

The types of actors and sources used are aligned with the framework: NATO/EU strategic materials, platform reports, disinformation analysis workshops (e.g., DFRLab, EU DisinfoLab, NATO StratCom COE), and the author's previous related works (Crimea/COVID/MI focus).

Theoretical Background: The Relationship Between FMK and Information Processing (OxIPO)

The methodological framework requires definitional discipline. If disinformation, "propaganda," cognitive warfare, and platform logic are mixed together, the analysis will be both debatable and irreproducible. This chapter is therefore

structured around Hungarian categories: all key terms are given in Hungarian, with their English equivalents in parentheses. The aim is not to engage in linguistic analysis, but to ensure that each chapter and each case (Crimea–COVID–Ukraine) can be read through the same conceptual lens.

Conceptual Frameworks and Hierarchy

In this study, weapon-grade communication (WGC) is defined as deliberately organized, strategic influence that attacks the target group's perception, decision-making, and behavior not only with statements, but also by shaping channels and environments. The essence of this is that the object of influence is not an opinion, but the processing chain itself: attention, credibility perception, meaning formation, and decision-making routines. Generative AI scales this effect: it not only produces more messages, but is also able to exploit ecosystem-level conditions (recommendation systems, moderation dynamics, comment ecosystems) as an operational space.

Consequence: a detection-focused approach alone is often insufficient, because the damage is not caused by a single false element, but by the overall effect (overload, thematization, erosion of trust).

For a disciplined comparison, I will describe the phenomenon on four inter-related levels:

- Information ecosystem: the technical, social, and regulatory environment (platform architecture, recommendation systems, moderation, media routines, user psychology) that determines the noise level and attention capacity load.
- Operation: a planned campaign with a defined time frame and goal that pursues a strategic objective (e.g., undermining trust, polarization, erosion of support).
- Tactic: a repetitive pattern of procedures that combines several techniques to achieve a goal. According to my own previous definition, FMK tactics are a pattern of procedures, not a single move: the coordinated use of several techniques in a situation/target group.
Example: the phenomenon of the "narrative arms race," in which alternative interpretations are built up in rapid succession, making it difficult to establish a stable framework.
- Technique: the specific means of executing a tactic (e.g., bot network, deep-fake, fake institutional identity, flood of comments). Technique is the "filling in" of the tactical level: the specific move that can be measured and detected.

The advantage of the hierarchy is that it does not allow technical observations of the "many bots" type to be confused with

operational intentions of the "wants polarization" type – and thus the indicators can also be recorded stably later on.

For the sake of conceptual order, the study treats FMK as an umbrella: FMK encompasses those operational patterns that cause damage (wrong decisions, erosion of public trust, institutional burden).

- Hybrid warfare: a strategic framework in which kinetic and non-kinetic means operate in synergy. Here, FMK is the element that serves cognitive dominance.
- Information operations: a doctrinal umbrella term; compared to this, FMK is a broader operational logic that is also interpreted in civilian platforms.
- Cognitive warfare: targets perception–interpretation–decision-making. The FMK organizes the "microtechniques" and "patterns" of its implementation into a measurable system.
- Infodemic: information overload, where processing capacity becomes a bottleneck. In relation to COVID-19, my own previous work also describes this state as a cognitive battlefield: biological spread is accompanied by information overload, which requires resilience.

The OxIPO Model as a Measurement Bridge

The methodological thesis of the study: security policy threats (FMK) can be

translated into information processing loads and errors. OxIPO is the "measurement bridge" for this. Ferenc Mező (2002), the developer of the OxIPO framework (Mező 2002, 2011, Mező & Mező, 2019), presents the OxIPO model as an organizational formula for human information processing, where performance is a function of organization and the Input–Process–Output chain.

In OxIPO, "organization" does not refer to an institution, but to a state of organization: how orderly the processing is, whether there is prior preparation, whether source criticism works, and whether there are fixed control points. We say "organization" instead of "organization" because the subject of measurement is not the organization, but the orderliness of the processing.

- High organization: preliminary schemas, control routines, stable interpretive framework → less manipulability.
- Low organization: ad hoc interpretation, emotional reactivity, attention fragmentation → greater overload.

FMK attacks the HIP chain on three points:

- Input: "occupation" of attention and perception (stimulus overload, rapid repetition, multi-channel dissemination). One well-known description of modern propaganda is the "firehose of falsehood," where the goal is often not persuasion, but overload and

disruption of the decision-making environment.

- Process: weakening deliberation through emotional tuning, framing, and uncertainty generation. Dowding and Oprea (political philosophy/political science, Economics & Philosophy) discuss political manipulation as something that undermines the conditions for rational deliberation—that is, it does not simply "persuade," but distorts decision-making autonomy.
- Output: flawed decisions, paralysis, delayed reactions, loss of trust.

"Degradation" (HIP degradation) is not a rhetorical turn of phrase here, but an operationalizable state:

- (a) overload – the speed/quantity of input exceeds processing capacity;
- (b) bias – emotional attunement overrides control;
- (c) latency – the decision gets stuck or falls apart.

Some modern forms of influence operate at the edge of consciousness (peripheral perception, accelerated stimuli, priming). Elgendi's (2018, *biomedicine/psychophysiology*; *Frontiers in Psychology*) review also emphasizes the methodological issues and limitations of subliminal (below-the-threshold) priming: the study of its effects is highly dependent

on masking, timing, and measurement protocols.

The study by Wang et al. (2021, cognitive neuroscience; *Brain and Cognition*) shows that in affective priming, emotional priming can measurably shift subsequent evaluation—that is, process-side bias can be empirically captured.

At the same time, Hernández-Gutiérrez et al. (2025, psycholinguistics; *Journal of Cognition*) urge caution: in many cases, “subliminal” linguistic priming also contains a conscious component, which is why the study treats this band not as “mystical mind control” but as a risk of loss of control.

Transition: the conceptual order of the chapter is critical because it is the only way to apply the same measurement logic to the Crimea–COVID–Ukraine arc in the later methodological sections (SFC + H–E–I/S–F–C + indicators).

Measurement Challenges and Literature Gaps

One of the most difficult yet unavoidable aspects of examining weaponized communication (WMC) is the treatment of intentionality. In the cognitive space, the “weapon” (narrative, meme, cloned news form, automated comment data) often behaves like an everyday mistake or part of a legitimate debate: the same sentence can be both a silly misunderstanding and an organized operational stimulus. However, if we completely remove intent from the analysis, FMK

loses its operational nature and slips back into describing a “noisy media environment.”

In the logic of OxIPO/HIP (Organization × [Input + Process + Output]), this is not a theoretical refinement: defense can only be activated in a targeted manner (especially in the Process stage) if the recipient—an individual, organization, or community—is able to separate accidental load (noise) from intentional load (operational signal). The aim of this chapter is therefore not to seek “internal intent” as a psychological state, but to derive it from external, observable patterns: in other words, to make intent examinable using proxy indicators.

Intentionality is a valid axis of analysis because types of misinformation (error, deception, malicious communication) are not merely “true–false” questions in practice, but rather interactions between intent and effect. The misinformation–disinformation–malinformation distinction is also based on this: it is not only the nature of the content but also the intention behind its production and dissemination that distinguishes the categories (Wardle–Derakhshan, 2017).

This study translates this into FMK language as follows:

- Unintentional error (misinformation): burdensome, but not “weaponized.” In HIP, it typically noises the Input, while the Process (verification) could in principle correct it.

- Intentional deception (disinformation): organized, repetitive, goal-oriented. In HIP, it targets the filters of Input and Process (authenticity, source, context) to force the Output onto a predetermined path.
- Real elements with harmful intent (malinformation): not necessarily false, but optimized for harm due to selection and framing. In HIP, it often shortens the Process stage by activating the emotional "fast track."

The second reason for intentionality is the normative boundary: the difference between legitimate persuasion and manipulation becomes clear where communication seeks to circumvent conscious deliberation or undermine deliberative control. Dowding and Oprea's (2024) theory of manipulation provides a useful framework here: "influence" becomes problematic not because it seeks to exert an effect, but because it reduces decision-making autonomy in the face of recognizability and justifiability (Dowding–Oprea, 2024). In FMK terms: the question is not whether it has an effect, but how it does so in the HIP chain.

Examining intent is tricky because "confessions" or clear chains of command are rarely available in the digital space. The solution, therefore, is not to "read minds," but to regulate intent attribution transparently: what external signs can we use to say, in a methodologically defensible way, that this is no longer a coincidence?

The study proposes four mutually reinforcing groups of indicators. Together, these constitute the "minimum for intent attribution" – the threshold above which we no longer speak of suspected FMK based on impression, but on methodological decision:

- Source identifiability: Organized deception typically works by hiding, masking, or cloning the source. If the source is systematically obscure (fake profile network, mirror site, "gray" intermediary channels), this increases the likelihood of intent. Aro's (2016) classic examples illustrate this well: coordination and disguise rarely occur by chance (Aro, 2016).
- Target audience identification: Accidental errors spread in a scattered manner. Intentional operations, on the other hand, resonate repeatedly with the same groups, with consistent emotional and identity connections. The target audience is not always declared, but the pattern of dissemination, linguistic "gate codes," and platform choice often indicate it.
- Damage pattern: We cannot measure internal intent, but we can measure the pattern of output consequences: does the content regularly undermine the conditions for orientation (trust, institutional cooperation, capacity for action), and is there a demonstrable stable trajectory of "paralysis–polarization–delegitimization"? This is where intentionality can be linked to HIP

deterioration: if the output persistently reinforces processing failures (delays, misacceptance, loss of sources), suspicion of FMK can be methodologically justified.

- Channel strategy: One of the strongest signs of organized operations is a multi-channel, repetitive, rhythmic appearance (text + image + video; multiple platforms; rapid domain switching). Prier (2017) also emphasizes from a military-strategic perspective that dominating trends and coordinating channels is not spontaneous "noise," but a predictable and executable pattern of behavior (Prier, 2017).

These four conditions are not "legal evidence" but a research decision rule: they are documentable, debatable, and therefore scientifically correct.

Éva Jakusné Harnos (2016), analyzing changes in methods of persuasion due to media development, points out that in linguistics, attribution of intent is a basic function, but in the online space, adaptability raises new questions. In classical mass communication (one-way broadcasting), intent can be inferred from editorial decisions. In the online space, however, content receives feedback (comments, shares, algorithmic ranking), which makes communication quasi-interactive: it reacts, fine-tunes, and repackages.

The methodological solution here is to treat intent as evidence of adaptation. If a narrative:

- quickly changes form (platform, language, visual solution) after countermeasures,
- maintains the same narrative arc in multiple modalities,
- and reorganizes the input tailored to the target audience, then the intention does not "disappear," but becomes visible in operation.

In OxIPO/HIP language: adaptivity indicates that the system does not merely produce content, but optimizes the way processing is circumvented—that is, the intention can be captured in the process pattern.

The intersection of weaponized communication (FMK) and human information processing (HIP) is not built from scratch. The literature offers plenty of measurement solutions, but they mostly run on separate tracks: cognitive psychology often measures micro-level mechanisms (stimulus, attention, priming) in the lab, while security policy and platform research practice tends to describe macro-level outputs (reach, network spread, coordination). Generative artificial intelligence (GenAI) and the multi-modal environment (text–image–sound–video) have brought these two levels together into a single operational space: stimulus flow has increased and industrial-scale distribution has become cheaper at the same time. Therefore, in this chapter, I provide a critical review of "previous experiments" and then identify the

research gap that the methodological framework of the study (OxIPO/HIP bridge + H-E-I grid + MRM) fills.

- Measuring cognitive load and attention capacity One classic way to measure damage to the HIP chain is to examine the load (how much stimulus enters, how much passes through processing): reaction time, error rate, memory accuracy, decision uncertainty. These measurements are good at capturing when reception shifts from "controlled" mode to "drifting" mode, but they often do not reveal whether the overload is the result of deliberate influence or mere information noise.
- Near-threshold stimulation and priming: One of the benefits of the debate surrounding subliminal/near-threshold effects is methodological discipline: how do we check whether the stimulus is perceived, how "visible" it is, and under what conditions we can talk about priming? Elgendi et al. (2018) review subliminal priming research not as a "mystical" phenomenon, but as a question of measurement protocols (e.g., masking, visibility tests, reaction time). The lesson for us is not that "anything can be done with invisible messages," but that in the age of FMK/GenAI, the perceived but uncontrolled stimulus zone is particularly important: where the recipient "lets through" the content peripherally because they do not have the capacity to activate control (Hernández-Gutiérrez et al., 2025).
- Proving delayed effects and lasting shifts: The study of hidden or semi-hidden learning is relevant to FMK because the goal is often not immediate persuasion, but rather the framing of later decisions. The experimental results of Ruch, Züst, and Henke (2016) show that certain stimulus associations can distort decisions even with a delay, while the participant is unable to explain the reason for the decision. The FMK interpretation of this is cautious but important: the operational effect often does not appear in "one big message" but in the sum of many small biases.
- Affective priming and emotional bias: Emotionally charged stimuli (fear, disgust, anger) often "take precedence" in the processing chain: they run faster and trigger automatic responses more easily. Wang et al. (2021) indicate that, in relation to affective priming, intense emotional facial expressions can activate different processing pathways than "normal" stimuli. This is important from an FMK perspective because, in a multimodal space (image + sound + caption), emotional overload reaches its target before the "cognitive brake" that triggers control.
- Measuring manipulation from a decision-making and normative perspective: Another way to measure political manipulation is not from the stimulus

side, but from the decision-making side: when rational reasoning is compromised, when "illegitimate" influence occurs. Dowding and Oprea (2024) approach manipulation from the perspective of, among other things, "evidential sincerity" and suitable reasons. From this, the study draws the methodological conclusion that "deterioration" in HIP is not just overload: it can be a deterioration in the quality of reasoning, i.e., when the recipient decides that they do not have access to the clear reasons necessary for their decision.

- Disinformation and network indicators (macro level): Security policy and platform-related research (e.g., NATO StratCom COE, EU/EEAS FIMI frameworks, fact-checking and disinformation labs) are strong in measuring dissemination and coordination: timing, channel synchronization, repetitive narrative panels, network acceleration, recurring "source washing." These indicators are excellent for describing the Situation (what is happening in space), but often weak for HIP-level response (what is happening in the recipient's processing), hence the need for a methodological bridge from OxIPO/HIP.

Based on a critical review, the *state of the art* remains incomplete in four areas compared to what FMK + GenAI + multimodality together represent.

- Integration gap (micro ↔ macro): Laboratory cognitive measurements (priming, load) and network/spread indicators rarely communicate with each other. There is no "translation key" that tells us what propagation pattern is likely to cause processing degradation in HIP. This study aims to connect this with the OxIPO/HIP organizing principle and the diagnostic logic of the H–E–I grid.
- Scale deficiency (genai-industrial scale): A significant portion of classical models implicitly assume human resource-limited operations. Generative AI, on the other hand, drastically reduces the cost and time of content production, so that the stimulus flow and number of variations are superimposed on the recipient as "synthetic noise." This cannot be treated as a mere "detector" issue: it creates a capacity and control problem in HIP.
- Multimodal deficiency (text-image-sound-video together): Measurements are often optimized for a single channel (text only, image only). Today's FMK operations, however, typically work in channel synchronization: the same frame appears in memes, videos, comment data, and "expert" articles. The research gap is therefore not whether "manipulation exists," but how control can be taught and measured across all major modalities. Lack of conceptual discipline in Hungarian: International terms are useful, but they

often "override" Hungarian logic. One of the aims of the study is to introduce key concepts in Hungarian in a logically clear manner and only then provide the English equivalent. This is not a linguistic game: a prerequisite for measurability is that the same phenomenon is always referred to by the same Hungarian category.

- Outcome: The above gaps together define the research gap: there is a need for an operationalizable methodology that treats (i) the operational nature of FMK, (ii) the scaling effect of GenAI, (iii) multimodal channel synchrony, and (iv) measurability as interpreted in HIP – all with Hungarian conceptual discipline. The following chapters therefore do not promise another case study, but rather a framework in which the phenomenon can be diagnosed, taught, and verified.

Methodological Proposal: The H–E–I Grid and Multimodal Resilience

After reviewing the previous chapter, the problem can be summed up in a single sentence: in the age of generative AI, it is not enough to recognize content; we need a framework that makes operational intent, propagation, and cognitive damage measurable. Accordingly, the aim of this chapter is to present an operationalizable (measurable–teachable–applicable) methodological set that combines qualitative comparison with indicator-based grids and a multimodal resilience toolkit.

The framework consists of three interrelated levels:

- (1) SFC as "background logic" (how we compare cases),
- (2) H–E–I grid as a diagnostic framework (what we ask in every case),
- (3) FMK matrix + RESILIENCE GRID + MRM as a "diagnosis to intervention" gateway (where we intervene and with what).

Structured-Focused Comparison (SFC) And Indicators

The methodological basis is provided by structured–focused comparison. According to the classic thesis of Alexander L. George and Andrew Bennett (political scientists, comparative case study methodology), the essence of SFC is that all cases examined are examined along the same set of questions ("structured") and narrowed down to the research problem ("focused") (George–Bennett, 2005). For me, SFC is not a "methodological embellishment" but a noise filter: the goal is to retain only those patterns from the endless sea of communication details that are relevant from an FMK perspective (armament, damage path, adaptivity).

Crimea 2014, COVID-19, and Ukraine 2022 are cases with different contexts, but the methodological study does not tell a story; rather, it answers the same questions using the same logic. The common set of questions can be divided into four blocks:

- Actor and intent: Who initiates the pattern, and what is the operational goal? (e.g., paralysis, undermining trust, rewriting legitimacy)
- Channel architecture: On what channels and with what channel synchronization does the operation run? (platforms, clone sites, network amplification)
- Narrative structure and emotional load (narrative & affect): What “cognitive shortcuts” does it build on? (fear, scapegoating, false certainty, urgency)
- Response time and resilience: How long does it take for the correction to appear, and how long does its effect last? (counter-narrative ratio, institutional-social response)
- Adaptivity: how quickly the channel or narrative of the operation changes in the event of a countermeasure (e.g., (T_{adapt}) , (T_{plat})).
- Amplification: the ratio of artificial acceleration and saturation (e.g., (Δ_{ampl})).
- Counter-narrative ratio: how visible and how “competitive” the correction is (e.g., (R_{enar})).

Indicators are not “smart words”: each indicator must specify whether it measures input, process, or output.

The Situation-Forecast-Control (SFC) Diagnostic Grid

Categories borrowed from international terminology are only helpful if they function according to Hungarian logical categories rather than as translations. Therefore, instead of S–F–C (Situation–Forecast–Control), I am establishing the following as the official Hungarian name:

The H–E–I grid is both a diagnostic and decision-making aid: it not only describes but also identifies points of intervention.

This set of questions can be immediately linked to the processing chain: input → process → output, which subsequent grids transform into a “measurement surface.”

I reinforce the qualitative logic of the SFC with a set of quasi-quantifiable indicators. The goal is not “perfect quantification,” but rather to be able to identify the same phenomena in the same way in three different cases.

The core of the identical set of indicators:

- Temporality: how much time elapses between the emergence of the narrative and effective correction (e.g., $(T_{\text{refutation}})$).

- *Situation: Initial Information Environment + Vulnerabilities.*

Here I record the initial noise level, the trust ecosystem, the target group's preliminary belief set, and the degree of “organization.” (Important clarification: “organization” here means not institution.)

- *Forecast: Expected Damage + Spread Logic.*

The forecast is not a prediction, but rather: through which channels and in what pattern the content is likely to spread, and what cognitive damage it is likely to cause (e.g., paralysis, panic, polarization, decision delay). As research on subliminal learning has confirmed, even unconscious stimuli can create relational connections in the hippocampus. In the H-E-I grid, this means that in the 'Forecast' phase, we must take into account not only conscious persuasion, but also delayed bias.

- *Control: Intervention Points + Resilience Tools.*

This is where the "what we do" side comes in: correction, exposure, channel-targeted intervention, institutional protocol, educational control. The goal is not to "erase" the narrative, but to mitigate its impact. Defense should not be limited to reinforcing the phenomenon of the "Liar's Dividend," i.e., not everything should be labeled "false" because that undermines trust in reality. The goal is specific, channel-based verification.

The H-E-I grid is OxIPO-compatible because it is able to keep the following together on a single sheet: (i) the input (Situation), (ii) the processing distortion (Prediction), and (iii) the output protection (Control).

The "engine" of the methodology is the combination of the attacking side and the

defending side: FMK matrix (which tactic-technique pairs work) + resilience grid (where and how we intervene). This is the point where diagnosis turns into intervention.

The reliability of the method is ensured by the fact that the case studies are analyzed using the same framework: I use the same set of H-E-I questions in all three vignettes and apply the same marking/coding rules to the indicators. This reduces the risk of one case being treated as a "story" and the other as a "measurement table": the comparison is thus not stylistic, but rule-based.

Validity in this framework is based on chain of evidence logic: every rating is backed by a traceable chain of sources (e.g., platform report / cybersecurity report / OSINT / fact-check / official announcement), and the framework also enforces negative testing. If the attribution, channel coordination, or level of organization cannot be substantiated, the rating is downgraded: I do not "prove" the FMK, but examine where the categorization does not hold.

The practical guarantee of falsifiability and reproducibility is that proxies and threshold logics are fixed: another researcher can produce ratings in a verifiable manner from the same secondary source base, or can dispute them (for example, by recalibrating thresholds or specifying alternative proxies).

This methodological discipline reduces the risk of anecdotal description and

closes the measurement-interpretation gap at the "PhD level."

In the FMK matrix, tactics are recurring patterns of behavior, while techniques are specific implementations. I record the following indicators for measurement:

- ($T_{\text{refutation}}$) (*Time To De-Bunk*): the time elapsed from appearance to the first widely accessible, credible correction. The "lingering" effect of misinformation has long been noted in cognitive psychology literature (e.g., continued influence), which is why delayed correction poses a particular risk (Ecker et al., 2015).
- (T_{adapt}) (*Adaptation Time*): the time it takes for the operation to change tactics/narrative in response to defensive measures.
- (T_{plat}) (*Platform Migration Speed*): how long it takes to migrate from the dominant channel to a new channel (e.g., "platform jump").
- (Δ_{ampl}) (*Amplification Delta*): the difference between artificial amplification and organic spread (on a proportional basis).
- (R_{enar}) (*Counter-Narrative Ratio*): the visibility/proportionality of the counter-narrative compared to the dominant narrative.

For qualitative comparability, I assign measurement proxies and scale values to each indicator. I determined the thresholds based on the reaction time analyses of NATO StratCom COE (2023) and the

memory consolidation windows described by Ecker et al. (2015). The thresholds are heuristic and aligned with typical reaction time windows in secondary sources; they can be recalibrated for primary data (Table 1). These indicators can typically be measured from several types of sources: platform reports, analyses by disinformation monitoring labs (e.g., DFR-Lab, EU DisinfoLab), institutional situation assessments, and, where possible, our own previous case study frameworks (Horváth, 2024).

Resilience is not "good intentions" but a tiered intervention system. The grid distinguishes between three levels (micro–meso–macro) and specifies the goal and the tool for each:

- *Micro (Individual)*: strengthening processing filters (attention management, emotional self-reflection, source criticism).
- *Meso (Organizational)*: procedures and protocols (crisis communication, control processes, "one source of truth" logic).
- *Macro (Social/State)*: regulation, institutional coordination, platform accountability, and strategic communication capabilities.

One of the stable tenets of resilience research is that "system resilience" does not consist of a single intervention, but is a layered, feedback-driven capability (Linkov–Trump, 2019).

Table 1: Measurement indicators and thresholds of the FMK matrix (Source: Author)

Indicator	Definition & Proxy	Scaling (Threshold Values)
\$T_{refutation}\$ (Response time)	Def: The time difference (\$\Delta t\$) between a disinformation narrative going viral and the first credible, wide-reaching institutional/fact-checking response.	Low (Good): < 4 hours (before consolidation).
		Medium: 4–24 hours (within the news cycle).
	Proxy: Google Trends / Crowdtangle spike vs. first Fact-Check URL timestamp.	High (Bad): > 24 hours (the narrative is fixed).
\$T_{adapt}\$ (Adaptivity)	Def: The attacker's reaction time to defensive measures (e.g., launching a new channel after being banned).	Low (Dangerous): < 6 hours (automated/pre-fabricated).
		Medium: 6–48 hours (reactive organization).
	Proxy: Domain registration dates and activity of new accounts after blocking.	High: > 48 hours (requires re-evaluation).
\$\Delta_{ampl}\$ (Amplification)	Def: The proportion of artificial (non-organic) distribution.	Low: < 10% (organic noise).
		Medium: 10–30% (supported campaign).
	Proxy: Proportion of CIB-suspicious accounts (creation date, posting frequency > 50/day) in total interaction.	High: > 30% (industrial-scale bot warfare).
\$R_{enar}\$ (Counterweight)	Def: The ratio of counter-narrative reach compared to the dominant FMK narrative.	Low: < 1:10 (the refutation is invisible).
		Medium: 1:10 – 1:2 (competitive situation).
	Proxy: Ratio of engagement (Like+Share+Comment) volumes for the two narratives.	High: > 1:2 (the counter-narrative dominates or is on par).

Multimodal Resilience Matrix (MRM)

In the era of generative AI, the threat is not only the quantitative increase in text content, but also the technical scaling of credibility hijacking. Previous research has shown that deepfake videos and voice-cloned content attack not rational judgment but "belief in evidence." Therefore, the defense cannot be uniform: the Multimodal Resilience Matrix (MRM) aims to assign specific verification routines to each modality (text, image, audio, video) in the process stage of the HIP chain.

Modalities and Control Points

The logic of the matrix is based on the micro–meso–macro resilience levels introduced in the study *"Defense Options,"* but here it is specifically broken down into sensory modalities (Table2).

TEXT – The field of scaled trust erosion: The main danger of generative language models (LLMs) is not individual lies, but the industrialization of "grammatically flawless" deception (e.g., *"Hi Mom, I'm in trouble"* type scams or *Doppelgänger* clone portals).

- Point of attack: Interpretation and logical inference. Flawless style disables previous "spelling suspicion" filters.
- Control: Stylistic Zero Trust. It is not linguistic correctness, but urgency and unusual channel switching that raise suspicion.

IMAGE – The illusion of proof: The *"Balenciaga pope"* or *"Pentagon explosion"* cases prove that photorealism activates reality heuristics in a fraction of a second, even before conscious cognition kicks in.

- Point of attack: First impression and emotional activation (limbic response).
- Control: Context validation. The image itself is not evidence, only an illustration. The question is not whether it "looks real," but whether there is an independent source for the event.

SOUND (AUDIO) – The most trusted gateway: Voice cloning (e.g., *Biden robocalls*, *CEO fraud*) is the most dangerous "last mile" attack, as familiarity with the tone of voice generates automatic trust (authority bias).

- Point of attack: Trust and personal identification. The voice "whispers in the ear," bypassing visual skepticism.
- Control: Callback protocol. The call should never be verified on the incoming channel, but on a known, authentic channel (another number, internal line).

VIDEO – A tool for paralysis: Although the *Zelensky* deepfake was not technically perfect, it demonstrated the "paralysis effect": moving images can temporarily paralyze decision-making with the power of visual shock.

- Point of attack: Perception of reality and illusion of presence.

- **Control:** Multi-source verification. Today, video is no longer “eyewitness” but constructed content. Authenticity is provided not by the visuals but by metadata and the chain of sources (Table2).

Operational Tool:

The "Stop-Check-Confirm" Routine

The operational, teachable element of MRM is the Stop–Check–Confirm routine, which works against the "delayed effect" described in cognitive psychology research. Since exposure to subliminal or manipulative stimuli leaves an immediate impression, the key to defense is to artificially slow down reaction time.

- *Stop:* If the stimulus (image/sound/text) triggers urgency, fear, or

euphoria, the primary reaction is a "cognitive emergency brake." This step eliminates the automatic emotional response.

- *Check:* What modality did it arrive in? (See matrix above). Is there a source, or is it just a "forwarded" message?
- *Confirm:* As described in the *Defense Options* study: is there independent confirmation on at least two separate channels? (E.g., the bank call in the app, the news from another news agency)

This routine does not require any technological tools, so it can form the basis of a "human firewall" in the cognitive domain of hybrid warfare.

Table 2: Multimodal Resilience Matrix (MRM): damage pathways and control points (Source: Author)

Modality	Attack Point	Typical Generative MI Threat	Resilience Control
TEXT	Interpretation and logical inference	massive comment and "article production", coherent false explanation	source triangulation, statement–evidence breakdown
IMAGE	first impression, emotional activation	synthetic photos used as "evidence"	image verification, context validation
AUDIO	trust and identification	voice cloning, imitation of "leadership" messages	callback protocol, authentic channel requirement
VIDEO	perception of reality, illusion of presence	deepfake, fake "eyewitness"	multi-source verification, metadata and context verification

Validation of the Framework: Case Studies (Triangulum)

The aim of this chapter is not to "narrate" three case studies, but to quickly test whether the methodological toolkit set out in earlier chapter (Situation–Forecast–Control, i.e. the S–F–C grid; as well as the FMK matrix and the logic of multimodal resilience) can be interpreted using the same set of questions to provide a comparable diagnosis and intervention picture in different types of conflicts. The "triangulum" therefore models three different cognitive load situations: (1) Crimea 2014 as a prototype (operational synchrony and input ambiguity), (2) COVID-19 as a laboratory/bridge (pure cognitive overload), (3) Ukraine 2022 as total application (technological mass production and the competition for resilience). The logic is the same in the OxIPO/HIP framework: input saturation and distortion, processing load, and then output and feedback control.

The framework uses the same set of H–E–I questions and the same basket of indicators for each vignette. In this chapter, I define eight briefly interpreted indicators: (1) O_organization (the target's preparedness and internal order), (2) V_vacuum (level of information vacuum/ambiguity), (3) S_saturation (saturation: quantitative and attentional overload), (4) M_multimodality (number and synchrony of modalities), (5) Δ_{ampl} (rate/sign of artificial amplification), (6)

T_refutation (delay of the first substantive, visible correction/refutation), (7) T_adapt (attacker's adaptation time to countermeasures), (8) R_enar (counter-narrative ratio/capacity). The goal is not to quantify everything, but to use the same measurement criteria throughout to ensure comparability (George–Bennett, 2005; Horváth, 2024).

Crimea 2014 (Prototype) – FMK And Operational Synchronization

Situation: In the Crimean operation, the communication space was not merely an accompanying phenomenon, but a framework that was timed to coincide with kinetic steps and paralyzed decision-making. "Operational synchrony" here means that physical events and the explanatory narratives about them converge in the same time window, so that the target receives "fact" and "interpretation" simultaneously, which reduces the chance of controlled processing (Holecz, 2018). In OxIPO/HIP terms, this can typically be described as low-medium O_organization and high V_vacuum: the rapidly changing situation, the uncertainty of identification, and the multiple, contradictory signals together increase the risk of paralysis.

Prediction? In the Crimean prototype, the damage path typically moves from uncertainty towards decision delay: if the input channels are ambivalent, the processing phase consumes too much energy

on basic situation interpretation, while capacity for answering the question "what should we do?" is reduced. The essence of reflexive control approaches in this framework is not to convince everyone, but to force the decision-maker and public opinion onto a collision course (Darczewska, 2014). Indicator image: S_saturation medium-high (rapid spread of a dominant framework), M_multimodality medium (classic media + local channels), Δ_{ampl} rather medium (limited amplification, partly based on human resources), T_adapt medium (narrative fine-tuning adapted to responses).

Control: Methodologically, the lesson of the prototype is that the first critical point of defense is not the debate over whether the content is true or false, but rather the timing of rapid, visible correction (T_refutation) and channel control. If the correction is delayed, false explanations can become permanently embedded and continue to influence decisions even after the refutation arrives (Ecker et al., 2015). In the Crimean case, T_refutation is typically high (on a scale of days/weeks), while R_enar is low- nd medium: defensive counter-narratives appear late and fragmented, so the controllability of feedback is limited.

Covid-19 (Labor/Híd) – Purely Cognitive Test Environment

Situation: The COVID-19 period functions as a "laboratory/bridge" case type

because the threat was largely cognitive, without a kinetic operational component, at both the global and local levels. The concept of infodemic describes precisely this state of saturation: the mixture of fragments of facts, fears, speculations, and deliberately spread misinformation causes processing overload and increases manipulability (Z. Karvalics, 2021). In the domestic context, the difficulty of defense is compounded by the fact that source criticism is often insufficient in the face of rapid, cross-platform dissemination (NMHH, 2021). Indicator image: S_saturation very high, V_vacuum medium (uncertainty surrounding the facts), O_organization highly dispersed (both at the individual and institutional levels), M_multimodality high (text–image–video together), Δ_{ampl} high (signs of algorithmic amplification and network dynamics).

Prediction: The damage path typically leads to exhaustion of the processing phase: too much conflicting information increases decision errors, unjustified certainty, or even cynical withdrawal. Methodologically, it is particularly important here to treat intentionality not as a state of mind, but as patterns and output effects (according to the logic of Chapter II). In the COVID lab, T_adapt is often low (rapid narrative shifts, new channels), while T_refutation is rather moderate (hours–days) because institutional corrections appear but do not immediately achieve a control effect due to saturation.

Control: Control here is not a single "big refutation," but a combination of resilience routines: proactive information, pre-bunking, and control measures tailored to different modalities (WHO, 2020; Horváth, 2025). Indicator image: R_enar is moderate, but often lags behind saturation; M_multimodality means that the effectiveness of defense is determined by whether it can provide control for each modality (e.g., visual control for visual misinformation, not just textual refutation).

Ukraine 2022– (Total Application) – Technological Mass Production and Resilience

Situation: After 2022, communication operations have become "industrialized" in several dimensions: cross-platform flow, synthetic content production, targeted network distribution, and rapid adaptation all appear together. At the same time, the target's O_organization is not the same as it was in 2014: war experience, memories of previous information attacks, and support from the international ecosystem often provide higher baseline resilience. The EU's reports on FIMI threats also emphasize that modern information manipulation is a systemic risk that requires ecosystem-level responses (EEAS, 2023).

Prediction: In total application, one of the keys to damage is T_adapt: whether the attacker network can respond quickly

to exposure, bans, domain loss, and platform sanctions. Doppelgänger-type operations are methodologically important precisely because they make the "return of intent" in patterns easy to grasp: when a network repackages itself in a matter of hours or days, moves to a new channel, and adopts a new narrative format, adaptation is a proxy for organization (EU DisinfoLab, 2023; EU DisinfoLab, 2024). Indicator image: T_adapt low, T_plat high (rapid platform migration), M_multimodality high (text–image–video–audio), Δ_{ampl} high (signs of network amplification and coordination). As the detailed ecosystem analysis of the 2022 war showed, the degree of Δ_{ampl} (amplification) was already on an industrial scale, in contrast to the human resource-intensive troll operations of 2014.

Control: Defense works here when T_refutation is low and feedback is rapid, not only institutionally but also communally (OSINT, fact-checking ecosystem). In addition, due to multimodality, resilience cannot be uniform: different controls are needed for text, images, audio, and video, especially when synthetic media appears (NATO StratCom COE, 2023; Horváth, 2025). Indicator image: T_refutation is low-medium (in many cases on an hourly scale), R_enar medium-high (strong counter-narrative capacity), but saturation (S_saturation) is still high at times, so the goal of control is not

to "reset" but to shorten the damage path and reduce decision errors.

Summary Validation: Cross-Sectional Image of the H–E–I Matrix

Table 3 compares the three cases (prototype, laboratory, total) based on the operationalized indicators recorded.

The purpose of the table is not primary measurement, but rather to organize the (secondary) proxies taken from the literature and OSINT analyses into a uniform grid, demonstrating that the framework is “workable” in practice.

Table 3: Comparison of the three case studies (Crimea, COVID, Ukraine) based on the H–E–I grid and indicators (Source: Author)

Case / Dimension	CRIMEA 2014 (<i>prototype</i>)	COVID-19 (<i>laboratory</i>)	UKRAINE 2022 (<i>total</i>)
Situation (Input) (Noise Level & Organization)	<i>High ambiguity</i> Low organization (Low O).	<i>Extreme noise (infodemic)</i> Changing organization	<i>High preparedness</i> High organization (High O).
Forecast (Process) (Expected damage trajectory)	<i>Paralysis</i> Indecisiveness due to uncertainty.	<i>Overload</i> Emotional/affective distortion.	<i>Attrition</i> (<i>exhaustion</i>) Erosion of trust and attention.
Tca'folat (Reaction time)	<i>High</i> (Days/Weeks)	<i>Moderate</i> (Hours/Days)	<i>Low</i> (< 4 hours, Pre- bunking)
Δ ampl (Amplification)	<i>Medium</i> (Trolls, manual)	<i>High</i> (Algorithmic echo chambers)	<i>Extreme</i> (GenAI bots, industrial scale)
Tadapt (Adaptivity)	<i>Medium</i> (Pre-made narratives)	<i>High</i> (Organic mutations)	<i>Low</i> (< 6 hours, automated)
Control (Output) (Resilience Focus)	<i>Reactive</i> Slow diplomatic and fact-finding response.	<i>Hybrid</i> Institutional (WHO) + Platform (tagging) response.	<i>Proactive</i> Pre-bunking, OSINT, multimodal control

Interpretation: Based on the table, alongside an increase in the attacking side's technological capacity (Δ_{ampl} , T_{adapt}), the defending side's reaction time ($T_{\text{refutation}}$) and organization also improve, which can be interpreted in line with the resilience-adaptation hypothesis.

Conclusions and Applications

The aim of the study was not to narrate a single conflict as a "story," but to establish an operationalizable framework in which weaponized communication (FMK) and generative artificial intelligence (GenAI) scaled influence patterns become comparable. The main conclusion is that defense in the cognitive space cannot be limited to refuting content after the fact: the quality of a decision is determined by the state and capacity of human information processing (HIP), which is why the target of intervention is often the processing chain itself (Input–Process–Output), rather than simply the classification of a statement as "true/false" (Ecker et al., 2015; NATO ACT, 2021).

The novelty of the study lies in four mutually reinforcing, directly usable "outputs." Together, these form the methodological package: definition → diagnosis → measurement → intervention.

1. Glossary hu→en (in parentheses), with definitional discipline. The glossary is not a translation aid, but a category record: it gives priority to the Hungarian conceptual order and only then provides

the international equivalents. The stakes are methodological: if the keywords are vague, the indicators will not be comparable either. Key elements (with examples):

- weapon-grade communication (WGC): organized, goal-oriented influence that distorts or overloads processing capacity and diverts decision-making.
- weapon-grade communication tactics: repetitive patterns of behavior within FMK (e.g., "fire hydrant effect").
- State of organization: the "O" element of OxIPO does not refer to an institution here, but to the state of preparedness of the system (resilience, protocols, preliminary schemes).
- Information ecosystem: the combined space of platforms, media, regulation, and trust.
- Saturation/overload, bias, paralysis: three practical "output" descriptors of HIP deterioration.

2. H–E–I grid (situation–forecast–control) as diagnostic logic. The H–E–I grid records the situation assessment as a framework for forecasting and intervention:

- situation: initial information environment, vulnerabilities, state of organization;
- forecast: expected damage trajectory and spread logic, probability of processing distortion;

- control: intervention points and resilience tools. Its novelty lies in the fact that it does not allow the analysis to slip between "description" and "recommendation": it keeps the diagnosis and response measures in the same coordinate system (George-Bennett, 2005; NATO ACT, 2021).
3. Multimodal Resilience Matrix (MRM). The MRM does not mix the modalities of text, image, sound, and video, but breaks them down into separate tracks and assigns control points to each. The significance of this becomes apparent in the era of generative AI: industrial scaling gives the attacking side not only quantitative advantages, but also the ability to switch modalities (Eliot, 2024; Wang et al., 2021).
4. Set of indicators (quasi-quantitative measurability). For the sake of comparability, the study establishes indicators that can be used to describe the dynamics of FMK between cases. Key indicators (for example): T_refutation (reaction time), T_adapt (adaptation time), T_plat (platform migration speed), Δ_{ampl} (artificial amplification), R_enar (counter-narrative ratio). The methodological benefit is that the question of intentionality can be approached not as a state of mind but as a pattern (proxy measurement): output regularity and adaptivity are among the strongest traces of organization (Bermal, 2021; Dowding-Oprea, 2024).

Limitations (Data Access, Source Criticism, Language Barriers)

Measurement sensitivity and data limitations.

The measurability of the proposed indicators is not homogeneous, which limits the depth of the analysis:

- Well measurable (OSINT-stable): T_refutation and R_enar can be estimated relatively stably from secondary sources based on public timestamps and reach/engagement data.
- Moderately measurable: Estimating T_adapt requires continuous monitoring, but can also be approximated from open sources (domain and account activity).
- Difficult to measure (requires platform data): Without internal platform data, the accurate determination of Δ_{ampl} can only rely on behavioral proxies, which increases the likelihood of false positives/false negatives.

Attribution and source criticism. Determining the "weapon-like nature" of FMK often depends on source identifiability, while concealment (gray/black communication) undermines methodological certainty. Here, the common minimum for defense and scientific correctness is confirmation from multiple sources and disciplined scaling of the strength of claims (Aro, 2016).

Linguistic and cultural barriers. A significant portion of the relevant material is in English/Russian/Ukrainian; even while

preserving Hungarian conceptual accuracy, there is a risk of mistranslation and loss of context. Including the domestic context (e.g., the vulnerabilities of the Hungarian media space, media literacy) reduces this distortion, but does not replace coverage of the entire linguistic spectrum (NMHH, 2021; Szicherle–Krekó, 2020).

Limitations of HIP-level impact measurement on a social scale. The measurement of HIP processes at the social level is typically indirect (behavioral data, public opinion, reach, response times). Therefore, the main strength of the framework is not "seeing into the mind," but comparing consistent patterns that indicate processing disorders (Csepeli, 1997; Ecker et al., 2015).

References

- Aro, J. (2016). The Cyberspace War: Propaganda and Trolling as Warfare Tools. *European View*, 15(1), 121–132. DOI <https://doi.org/10.1007/s12290-016-0395-5>
- Atlantic Council – DFRLab (2024). AI tools usage for disinformation in the war in Ukraine. *Atlantic Council DFRLab*. [online].
- Bermal, P. (2021). Misinformation: Why Most Fake News Regulation is Doomed to Failure. *British Association of Comparative Law*, January 29, 2021. [online].
- Birds Aren't Real (2021). *Birds Aren't Real - The Movement*. [online]. Available at: birdsarentreal.com
- Brookie, Graham et al. (2021). *Weaponized: How Rumors About COVID-19 Origins Led to a Narrative Arms Race*. Washington, DC: Atlantic Council DFRLab.
- CNA (2023). *Assessing Russian Cyber and Information Warfare in Ukraine: Expectations vs. Reality*. Arlington: CNA. [online].
- Cresswell, K. M., Worth, A., and Sheikh, A. (2010). Actor-Network Theory and Its Role In Understanding the Implementation of Information Technology Developments in Healthcare. *BMC Medical Informatics and Decision Making*, 10, 67. DOI <https://doi.org/10.1186/1472-6947-10-67>
- Csepeli, Gy. (1997). *Social Psychology*. Budapest: Osiris.
- Darczewska, J. (2014). *The Anatomy of Russian Information Warfare: The Crimean Operation, A Case Study*. Warsaw: OSW. [online].
- Dowding, K. and Oprea, A. (2024). Manipulation in politics and public policy. *Economics and Philosophy*, 40(3), 685–710. DOI <https://doi.org/10.1017/S0266267124000063>
- Ecker, U. K. H. et al. (2015). He did it! She did it! No, she did not! Multiple causal explanations and the continued influence of misinformation. *Journal of Memory and Language*, 85, 101–115. DOI <https://doi.org/10.1016/j.jml.2015.06.003>
- Elgendi, M. et al. (2018). Subliminal Priming—State of the Art and Future Perspectives. *Behavioral Sciences*, 8(6), 54. DOI <https://doi.org/10.3390/bs8060054>
- Eliot, L. B. (2024). Generative AI Uses Subliminal Messaging For Shadowy Persuasion. *Forbes*, October 20, 2024. [online]. Available at: [Forbes.com](https://www.forbes.com)
- EU DisinfoLab (2024). *The Doppelgänger Operation*. Brussels: EU DisinfoLab. [online].
- European External Action Service (EEAS) (2023). *2022 Report on Foreign Information Manipulation and Interference (FIMI) Threats*.

- Brussels: EU External Action Service. [online]. Available at: ec.europa.eu
- Galeotti, M. (2018). I'm Sorry for Creating the 'Gerasimov Doctrine'. *Foreign Policy*, March 5, 2018. [online].
- George, A. L. and Bennett, A. (2005). *Case Studies and Theory Development in the Social Sciences*. Cambridge: MIT Press.
- Giles, K. (2016). *Russia's 'New' Tools for Confronting the West*. London: Chatham House.
- Hargitai, D. and Józsa, L. (2012). The relationship between organizational resilience and leadership. *Management Science*, 43(Ksz), 48–58.
- Hernández-Gutiérrez, D. et al. (2025). The Conscious Side of 'Subliminal' Linguistic Priming: A Systematic Review With Meta-Analysis and Reliability Analysis of Visibility Measures. *Journal of Cognition*, 8(1), 13. DOI <https://doi.org/10.5334/joc.419>
- Holecz, J. (2018). Chronology of the 2014 occupation of the Crimean Peninsula. *Military Science Review*, 11(3), 68–85.
- Horváth, D. (2024). Fake news, or communication tactics classified as weapons, as a current challenge in military science. In Gazdag, F., Padányi, J. and Tálas, P. (Eds.). *Current Issues in Military Science 2023*. Budapest: Ludovika University Press. pp. 151–168.
- Jakusné Harnos, É. (2016). Changes in Methods of Persuasion in the Light of the Development of the Media. *Honvédségi Szemle: The Central Journal of the Hungarian Armed Forces*, 144(1), 180–190.
- Jowett, Garth S. and O'Donnell, Victoria (2011). *Propaganda & Persuasion*. Thousand Oaks: SAGE Publications.
- Kofman, M.I and Rojansky, M. (2015). A Closer Look at Russia's 'Hybrid War'. *Kennan Cable*, No. 7. Wilson Center.
- Linkov, I. and Trump, B. D. (Eds.) (2019). *The Science and Practice of Resilience*. Cham: Springer.
- Locoman, O. (2024). Russian state TV coverage of Ukraine 2013–2014 and 2021–2022. *ResearchGate*.
- Media1 (2020). Traffic to Hungarian news sites has skyrocketed due to the coronavirus pandemic. *Media1.hu*, March 29, 2020 [online].
- Mező F. (2002). *A tanulás stratégiája*. Pedellus Novitas Kft., Debrecen
- Mező F. (2011). *Tanulás: diagnosztika és fejlesztés az IPOO modell alapján*. K+F Stúdió Kft., Debrecen.
- Mező F., & Mező K. (2019). Az OxIPO-modell – az interdiszciplináris kutatások egy lehetséges értelmezési kerete. *OxIPO – interdiszciplináris tudományos folyóirat*, 2019/1, 9–21. DOI <https://doi.org/10.35405/OXIPO.2019.1.9>
- National Media and Infocommunications Authority (NMHH) (2021). We would defend ourselves against fake news, but this is not always successful – NMHH research. *NMHH.hu*, October 25, 2021. [online]. Available at: nmhh.hu
- National Media and Infocommunications Authority (NMHH) (2022). The fight against fake news is intensifying: we cannot stay out of it. *NMHH.hu*, March 24, 2022 [online].
- NATO Allied Command Transformation (ACT) (2021). *Cognitive Warfare Concept Note*. Norfolk: NATO ACT Innovation Hub. [online].
- NATO StratCom COE (2015). *Analysis of Russia's Information Campaign Against*

- Ukraine*. Riga: NATO StratCom COE. [online].
- NATO Strategic Communications Centre of Excellence (NATO StratCom COE) (2023). *Virtual Manipulation Brief (VMB) 3/1*. Riga: NATO StratCom COE. [online]. Available at: stratcomcoe.org
- Naughton, J. (2020). Fake News about Covid-19 Can Be as Dangerous as the Virus. *The Guardian*, March 14, 2020. [online].
- NLC.hu (2017). Balaton Deniers Club. *NLC.hu*, August 9, 2017. [online].
- Nolan, S. and Kimball, M. (2021). The Intent Behind a Lie: Mis-, Dis-, and Malinformation. *Psychology Today*, December 2021. [online]. Available at: psychologytoday.com
- Prier, J. (2017). Commanding the Trend: Social Media as Information Warfare. *Strategic Studies Quarterly*, 11(4), 50–85. [online].
- Rothkopf, David (2003). When the Buzz Bites Back. *The Washington Post*, May 11, 2003. [online].
- Ruch, S., Herbert, E., and Henke, K. (2017). Subliminally and sup-raliminally acquired long-term memories jointly bias delayed decisions. *Frontiers in Psychology*, 8, 1542. DOI: <https://doi.org/10.3389/fpsyg.2017.01542>
- Ruch, S., Züst, Marc Alain, and Henke, K. (2016). Subliminal messages exert long-term effects on decision-making. *Neuroscience of Consciousness*, 2016(1), niw013. DOI <https://doi.org/10.1093/nc/niw013>
- Salam, E. (2021). Majority of Covid Misinformation Came from 12 People, Report Finds. *The Guardian*, July 17, 2021. [online].
- Síklaki, I. (1997). *The Psychology of Persuasion*. Budapest: Scientia Humana.
- Sun T. (1995). *The Art of War*. (Trans. Ferenc Tókei). Budapest: Balassi.
- Szenes, Z. (2021). Policy against hybrid threats in Hungary. *Military Science*, 31(E-issue), 5–22.
- Szicherle, P. and Krekó, P. (2020). Fake news proliferates in the wake of the coronavirus. *Political Capital*, March 16, 2020 [online].
- TalTech (2022). *The Evolution of Russian Propaganda in Ukraine (2014→2022)*. Tallinn: TalTech.
- Uscinski, J. E. and Parent, J. M. (2014). *American Conspiracy Theories*. New York: Oxford University Press.
- Wang, Y. et al. (2021). Subliminal affective priming effect: Dissociated processes for intense versus normal facial expressions. *Brain and Cognition*, 148, 105674. DOI: <https://doi.org/10.1016/j.bandc.2020.105674>
- Wardle, C. and Derakhshan, H. (2017). *Information Disorder: Toward an interdisciplinary framework for research and policy making*. Strasbourg: Council of Europe. [online].
- WHO (2020). *COVID-19 Vaccine Safety Communication*. World Health Organization. [online].
- Z. Karvalics, L. (2021). Infodemic. *UNESCO Hungarian National Commission, Information for All Program*. [online]. Available at: unesco.hu