

DOKTORI (Ph.D) ÉRTEKEZÉS

Szabadföldi István

**A Mesterséges Intelligencia alkalmazási lehetőségei a
NATO Átfogó Művelettervezési Direktíva (COPD)
módszertana alapján történő katonai művelettervezés
támogatásában**

— 2026 —

NEMZETI KÖZSZOLGÁLATI EGYETEM

Katonai Műszaki Doktori Iskola

Témavezető:

Dr. habil. Négyesi Imre Ph.D

Budapest, 2026.





TARTALOMJEGYZÉK

1. BEVEZETÉS.....	12
1.1. A kutatás időszerűsége és relevanciája (NATO/EU/MH kontextus)	13
1.2. Tudományos problémafelvetés és kutatási rés (research gap).....	14
1.3. Kutatási kérdések	15
1.4. Hipotézisek	15
1.5. Kutatási célok és célkitűzések.....	16
1.6. Alkalmazott módszerek és adatforrások.....	16
1.7. Lehatárolások, feltételezések, korlátok	17
1.8. Az értekezés felépítése és logikája	18
2. TAXONÓMIAI KERET – A MESTERSÉGES INTELLIGENCIA ÉS FELTÖREKVŐ/DISZRUPTÍV TECHNOLÓGIÁK A KATONAI KÖRNYEZETBEN	19
2.1. Feltörekvő és diszruptív technológiák (EDT) – fogalom és NATO-értelmezés.....	19
2.1.1. A NATO EDT-konceptió intézményi beágyazottsága és politikai–katonai célrendszere	20
2.1.2. EDT-k és a műveleti tervezés: releváns hatásmechanizmusok.....	21
2.2. MI alapfogalmak és taxonómia (ML/DL, generatív MI, hibrid rendszerek)	23
2.2.1. Gépi tanulás (ML): felügyelt, felügyelet nélküli és megerősítéses tanulás.....	25
2.2.2. Mélytanulás (DL) és alapmodellek: a generatív MI katonai jelentősége	27
2.2.3. Szimbolikus és hibrid (neuro-szimbolikus) rendszerek: miért relevánsak a tervezésben?..	29
2.3. Katonai alkalmazási területek (ISR, C2, EW, logisztika, cyber).....	29
2.3.1. ISR és OSINT: össz-adatforrású feldolgozás és fúzió	30
2.3.2. C2 és tervezéstámogatás: COA-generálás, wargaming, optimalizáció	33
2.3.3. Elektronikai hadviselés (EW): spektrum, jelazonosítás, adaptív zavarás.....	34
2.3.4. Logisztika és fenntartás: predikció, ellátási hálózatok és erőforrás-optimalizáció.....	35
2.3.5. Kiber (Cyber): anomáliadetektálás, fenyegetés-intelligencia, aktív védelem.....	35
2.3.6. Összegző megállapítások	36
2.4. Megbízhatóság, magyarázhatóság, adatminőség és kiberbiztonsági követelmények	36
2.4.1. Megbízhatóság és robusztusság: „pontosság” helyett életciklus-szemlélet	37
2.4.2. Magyarázhatóság, nyomon követhetőség és auditálhatóság	38
2.4.3. Adatminőség és adatgazdálkodás: a katonai MI szűk keresztmetszete	39
2.4.4. Kiberbiztonság és adversarial fenyegetések: MI mint támadási felület	40
2.4.5. VVTE, minőségbiztosítás és folyamatos monitorozás (MLOps)	41
2.4.6. Követelményrendszer-összegzés: minimum-követelmények katonai MI-hez	42
2.5. MI és autonómia kapcsolata; emberi kontroll (human-in-the-loop/on-the-loop)	43
2.5.1. Fogalmi keret: autonóm, fél-autonóm és operátor-felügyelt rendszerek.....	43

2.5.2. Emberi kontroll: felelősség, beavatkozási jog és „meaningful human control”	44
2.5.3. Autonómia, etikusság és jogi megfelelés: következmények a tervezéstámogatásra	44
2.6. Összefoglalás	45
3. A KATONAI MŰVELETTERVEZÉS FEJLŐDÉSE ÉS ELMÉLETI ALAPJAI	46
3.1. A művelettervezés fogalmi rendszere (stratégiai–hadműveleti–harcászati szint)	46
3.2. Történeti áttekintés: ókortól a napóleoni háborúig	48
3.3. Ipari korszak és világháborúk: a modern hadműveleti művészet kialakulása	49
3.4. Hidegháború és poszthidegháború: joint, expeditionary, hálózatalapú műveletek.....	49
3.5. Multi-domain és információs kor: komplex adaptív műveleti környezet	50
3.6. Összefoglalás	50
4. NATO ÉS EU MŰVELETTERVEZÉSI DOKTRÍNÁK ÉS MÓDSZERTANOK.....	52
4.1. Átfogó megközelítés (comprehensive approach) és a DIME/PMESII keretrendszerek	52
4.2. NATO AJP-5: a hadművelettervezés doktrínája	54
4.3. NATO COPD v3.1: cél, helye és kapcsolódása az AJP-5-höz.....	56
4.4. EU katonai tervezés: CONOPS–OPLAN logika, MPCC/EUMS szerepek.....	58
4.5. Magyar katonai doktrinák és szabályzók kapcsolódásai	60
4.6. Összefoglalás	61
5. A NATO COMPREHENSIVE OPERATIONAL PLANNING DIRECTIVE (COPD) V3.1 MÓDSZERTAN RÉSZLETES ELEMZÉSE	62
5.1. COPD alapelvek, szerepkörök, tervezési termékek rendszere.....	63
5.1.1. Tervezési szervezet: JOPG, working groupok és boardok.....	64
5.1.2. Battle rhythm, „quality gates” és a tervezési minimum fogalma.	64
5.1.3. Tervezési termékek rendszere: SSA–MRO–CONOPS–OPLAN.....	65
5.1.4. Traceability és auditálhatóság.	65
5.1.5. Információmegosztás, besorolás és „releasability” rétegezés.	67
5.1.6. Tudásmenedzsment, verziókezelés és változásnapló.....	67
5.1.7. Döntési napló és feltételezés-menedzsment.....	68
5.2. Fázisok és főtermékek (Phase 1–6) – döntési pontok és jóváhagyások	68
5.2.1. Phase 1 – Initial Situational Awareness (ISA) és tervezésindítás	69
5.2.2. Phase 2 – Strategic Assessment (SSA): küldetés- és környezet-elemzés	70
5.2.3. Phase 3 – MRO/COA-fejlesztés: alternatívák létrehozása és tesztelése.....	71
5.2.4. Phase 4 – CONOPS és OPLAN/OPORD: integrált részletezés és végrehajthatóság	72
5.2.5. Phase 5 – Végrehajtás, assessment és FRAGO (adaptáció)	73
5.2.6. Phase 6 – Transition/Termination: kilépés, átadás és szervezeti tanulás.....	74
5.2.7. Gyorsított tervezés (accelerated planning) – a sebesség és bizonytalanság kezelése	74
5.3. Műveletkialakítás (operational design): COG, LoO/LoE, hatások, kockázat	76

5.3.1. COG-elemzés és a CG–CC–CR–CV lánc gyakorlati értelme	76
5.3.2. Decisive conditions (DC) és objectives (OOs).....	77
5.3.3. LoO/LoE és szinkronizáció.....	77
5.3.4. Hatásláncok (effects tree) és mérhetőség.....	77
5.3.5. Kockázatkezelés: risk register, risk appetite és mitigáció	78
5.4. COA-fejlesztés, wargaming és összehasonlító értékelés	79
5.4.1. Wargaming és red teaming.....	79
5.4.2. COA-összehasonlítás és döntési transzparencia.....	79
5.4.3. Decision Support Template/Matrix.	80
5.5. Műveletértékelés (assessment), mérőszámok, lessons learned beépítése	81
5.5.1. MOP és MOE, leading és lagging indikátorok.	81
5.5.2. Data governance és besorolási korlátok	82
5.5.3. Lessons learned intézményesítése.....	82
5.6. Jelen kihívások a NATO hadműveleti tervezésben (időnyomás, adatbőség, komplexitás)	83
5.6.1. Időnyomás és accelerated planning.....	83
5.6.2. Adatbőség, OSINT és dezinformáció.....	83
5.6.3. Multidomain komplexitás és interoperabilitás.....	83
5.7. Összefoglalás.....	84
6. FEJEZET MI TÁMOGATÁS A COPD 1-2 FÁZISOKBAN: A MŰVELETI KÖRNYEZET MEGÉRTÉSE, STRATÉGIAI HELYZETÉRTÉKELÉS (INITIAL SITUATIONAL AWARENESS, STRATEGIC ASSESSMENT)	85
6.1. A DIME/PMESII és kapcsolódó elemzési keretek (ASCOPE/ICR2) – adatmodell és indikátorok, a PMESII, JIPOE és össz-adatforrású felderítés adatforrásai	85
6.1.1. Adatmodell: entitások, kapcsolatok, események és indikátorok	90
6.1.2. DIME–PMESII összerendelés és indikátor-példatár	91
6.1.3. Adatforrások: PMESII, JIPOE és össz-adatforrású felderítés.....	92
6.2. JIPOE folyamata és termékei; kapcsolódási pontok a COPD-hoz	93
6.2.1. A JIPOE lépések és a COPD 1-2 fázisok összekapcsolása	93
6.3. Adatgyűjtési módszerek és adatforrások: HUMINT, SIGINT, GEOINT/IMINT, OSINT, CYBINT ..	94
6.3.1. HUMINT.....	95
6.3.2. SIGINT.....	95
6.3.3. GEOINT/IMINT	96
6.3.4. OSINT.....	96
6.3.5. CYBINT (kibertéri információk és fenyegetési hírszerzés)	97
6.3.6. Összefoglaló: források, adatfajták és MI-támogatási pontok	97
6.4. Adatfúzió és elemző lánc (intelligence cycle); dezinformáció és deception kezelése	98
6.4.1. Adatfúziós architektúra: a nyers adatoktól a döntéstámogató termékig.....	99

6.4.2. Dezinformáció és deception: fenyegetések és kezelési módszerek	99
6.4.3. MI-kockázatok és kontrollok az elemző láncban	100
6.5. Fenyegetésvértékelés (Threat Assessment) prediktív elemzése	101
6.5.1. Prediktív módszerek a fenyegetésvértékelésben	102
6.5.2. Minőségbiztosítás és ember–gép együttműködés	102
6.6. MI eszközök az elemzés támogatására (NLP, CV, gráf- és hálózatelemzés, anomáliaészlelés)	103
6.6.1. Természetesnyelv-feldolgozás (NLP)	103
6.6.2. Számítógépes látás (Computer Vision, CV) és térképi analitika	104
6.6.3. Gráf- és hálózatelemzés	105
6.6.4. Anomáliaészlelés és riasztás (alerting)	105
6.6.5. Integráció a törzsmunkába és termékekbe (COP, SSA, briefek)	106
6.7. Részletes módszertani kiegészítések és megvalósítási javaslatok	107
6.7.1. Indikátor-katalógus kialakítása (PMESII-PT/ASCOPE szemlélet)	107
6.7.2. JIPOE termékek részletesen és MI-vel támogatható előállításuk	108
6.7.3. Gyűjtési követelmények és collection management a COPD 1-2 fázisában	109
6.7.4. Elemzői tradecraft: strukturált technikák, torzítások és megbízhatóság	110
6.7.5. Prediktív modellek validációja és mérőszámok	111
6.7.6. Bevezetési útiterv és szervezeti feltételek	111
6.9. Összegzés	112
7. FEJEZET MI TÁMOGATÁS A COPD 3-4 FÁZISÁBAN: KATONAI VÁLASZ OPCIÓK, KONCEPCIÓ ÉS TERVEZÉS (MILITARY RESPONSE OPTIONS, CONOPS, OPLAN)	114
7.1. Mesterséges intelligencia által vezérelt cselekvési terv (MRO/COA) generálása és szimulációja	115
7.1.1. Bemenetek, korlátok és tervezési „szerződés”: mit kell a MI-nek tudnia?	116
7.1.2. COA/CONOPS adatmodell: feladatok, hatások, erőforrások és idő–tér összerendelése	117
7.1.3. COA-generálási megközelítések: szabály, eset, keresés és generatív MI	119
7.1.4. COA-elemzés szimulációval: wargaming, agent-based és sztochasztikus értékelés	121
7.1.5. Kimenetek: COA-brief, összehasonlító mátrix és „tervezési emlékezet”	122
7.1.6. Átmenet a COA-tól a CONOPS-ig: MI a koncepció kidolgozásában	123
7.1.7. OPLAN dokumentum-összeállítás és mellékletek: sablonkitöltés, konzisztencia és biztonság	124
7.2. Optimalizáló algoritmusok erőforrás-elosztáshoz	125
7.2.1. Optimalizációs problémátípusok a tervezésben	126
7.2.2. Determinisztikus, robusztus és sztochasztikus optimalizáció	126
7.2.3. Heurisztikák, metaheurisztikák és megerősítő tanulás	126
7.2.4. Optimalizáció és MI-kormányzás: átláthatóság, audit és felelős használat	127

7.2.5. Többcélú optimalizáció és Pareto-szemlélet a COA-választásban.....	128
7.2.6. Erőforrás-elosztás különleges területei: ISR, információs műveletek és cyber	128
7.2.7. Érzékenységvizsgálat, „what-if” elemzés és red teaming	129
7.2.8. Adatminőség és interoperabilitás: az optimalizáció „üzemeltethetősége”	129
7.3. Integráció háborús játékok eszközeivel	130
7.3.1. MI a wargaming előkészítésében: scenárió, eseménykártya, ellenség-modellezés	130
7.3.2. MI a wargaming levezetésében: félautomata adjudikáció és eseménynapló	130
7.3.3. M&S eszközök és MI: adatcsere, digitális iker és kimeneti analitika	131
7.3.4. Szervezeti és emberi tényezők: bizalom, torzítás és standardok	131
7.3.5. Információbiztonság és adatkezelés a wargaming–MI integrációban.....	132
7.4. Empirikus elemzés: prototípusok és szimulációk.....	133
7.4.1. Kutatási megközelítés és prototípus-logika	133
7.4.2. Prototípus-architektúra: COA-generátor, optimalizáló modul és szimulációs hurok.....	134
7.4.3. Reprodukálható szimulációs példa: két COA összehasonlítása bizonytalanság mellett...	134
7.4.4. Erőforrás-optimalizáció demonstráció: VOI-alapú ISR-kiosztás.....	136
7.4.5. Prototípusok értékelése: javasolt mérőszámok és kontrollók.....	137
7.4.6. Összegző következtetések: mit ad hozzá a MI a COPD 3–4. fázisban, és hol vannak a korlátok?	137
8. FEJEZET MI TÁMOGATÁS A COPD 5-6 FÁZISÁBAN: VÉGREHAJTÁS, ÁTMENET (EXECUTION, TRANSITION)	139
8.1 MI a tervvalidációhoz és az adaptív újratervezéshez.....	139
8.1.1 Tervvalidáció: megfelelőség, konzisztencia és végrehajthatóság.....	140
8.1.2 MI-alapú validációs eszköztár: szabályok, szimulációk és „digitális iker” megközelítés....	142
8.1.3 Adaptív újratervezés: döntési pontok, ágak és folytatások MI-támogatással	143
8.2 Valós idejű monitorozás és visszacsatolási hurkok.....	145
8.2.1 Adatarchitektúra a végrehajtásban: adatfolyamok, metaadatok és döntési ritmus.....	146
8.2.2 Visszacsatolási hurkok: operatív értékelés és „zárt láncú” döntéstámogatás.....	148
8.2.3 Dezinformáció, deception és adversarial ML a végrehajtási adatokban	149
8.3 Művelet utáni elemzés MI segítségével.....	150
8.3.1 AAR adatforrások, adatminőség és bizonytalanság.....	151
8.3.2 Automatizált feldolgozás: NLP, multimodális elemzés és hálózati megközelítés	152
8.3.3 Tudásmenedzsment és átmenet: a tanulságok visszacsatolása a következő ciklusba	153
8.4 Megvalósítási kihívások (pl. adatbiztonság, interoperabilitás).....	154
8.4.1 Adatbiztonság, kiberreziliencia és ellátási lánc.....	154
8.4.2 Interoperabilitás és adatstandardok a szövetségi végrehajtásban	155
8.4.3 Kormányzás, felelős MI és jogi megfelelőség	157

8.4.4 Szervezeti, képzési és bevezetési kihívások: képességfejlesztés a Tranzíciós fázis tükrében	158
8.5 Összegzés	160
9. FEJEZET - KOCKÁZATOK, ETIKAI ÉS JOGI MEGFONTOLÁSOK, KIBERBIZTONSÁG	161
9.1. Felelős MI (Responsible AI) és emberi kontroll elvei	162
9.1.1. RAI elvek operationalizálása a katonai MI-életciklusban.....	164
9.1.2. Emberi kontroll a döntési ciklusban: automatizációs torzítás és kalibrált bizalom	165
9.2. Nemzetközi humanitárius jog, NATO/EU/USA policy-k és standardok	166
9.2.1. Policy-keretek közös nevezői és feszültségpontjai	167
9.2.2. A „kritikus funkciók” és a felelősség problémája.....	168
9.3. Adatvédelem, torzítás, explainability, félreinformálás és manipulációs kockázatok	169
9.3.1. Adatvédelem és adatbiztonság: minősített és nyílt adatok együttélése	170
9.3.2. Torzítás, hibamódok és magyarázhatóság a törzsmunkában	171
9.3.3. Manipuláció és dezinformáció: kockázati indikátorok a műveleti környezetben	171
9.4. Adversarial ML, modellmérgezés, ellentevékenységek; védekezési stratégiák	172
9.4.1. AML fenyegetési térkép: támadási felületek és védelmi kontrollok	174
9.4.2. Kiberbiztonság és MI: integrált incidenskezelés és tanuló védelem	174
9.5. Minőségbiztosítás, tesztelés-értékelés (T&E), megfelelés és audit.....	175
9.5.1. T&E ellenőrzőlista katonai döntéstámogató MI-re.....	176
9.5.2. A megfelelés bizonyítása: dokumentáció, naplózás, „evidence by design”	177
9.6. Összefoglalás.....	177
10. ÖSSZEGZETT KÖVETKEZTETÉSEK.....	179
10.1. Az elvégzett vizsgálat tömör leírása és fejezetenként összegzett következtetésem.....	179
10.1.1. Elméleti keret: a tervezéstámogató MI minimum-követelményei.....	180
10.1.2. A katonai művelettervezés fejlődése: miért időszerű az MI-integráció?	180
10.1.3. NATO és EU tervezési doktrínák: közös pontok és követelmények az MI számára.....	181
10.1.4. COPD v3.1 elemzés: MI-beavatkozási pontok a tervezési termékek mentén	182
10.1.5. COPD 1–2. fázis: MI a műveleti környezet megértésében és a stratégiai helyzetértékelésben.....	182
10.1.6. COPD 3–4. fázis: MI a COA/MRO, CONOPS és OPLAN kidolgozásában	183
10.1.7. COPD 5–6. fázis: MI a végrehajtásban, az átmenetben és az adaptív újratervezésben .	183
10.1.8. Kockázatok, etika-jog és kiberbiztonság: az MI-támogatás fenntarthatósági feltételei.	184
10.1.9. Integrált következtetés: az MI-támogatás értéke és korlátai a COPD egészében	185
10.2. Új tudományos eredményeim	186
10.3. Gyakorlati felhasználási ajánlások	188
10.3.1. Az AI technológiák alkalmazása a katonai művelettervezésben.....	188

10.3.2. MH/NATO – eljárásrend, képességfejlesztés, bevezetési útiterv	189
10.3.3. Interoperabilitás növelése a NATO erőkön belül	191
10.3.4. Módszertani ajánlások a mesterséges intelligencia hatékony alkalmazásához	192
10.3.5. A COPD, a PMESII adatforrások és az MI integrációja INCOSE Systems Engineering szemlélettel	193
10.4. Mérőszámok a mesterséges intelligencia katonai művelettervezésben történő felhasználásának értékeléséhez	194
10.5. Ajánlás további kutatási irányokra	196
IRODALOMJEGYZÉK	198
PUBLIKÁCIÓS JEGYZÉK	206
MELLÉKLETEK	208
Melléklet A – Tervezési sablonok és kitölthető űrlapok (COPD-kompatibilis)	208
A.1. Döntési napló (Decision Log) – sablon	208
A.2. Feltételezéslista (Assumption Register) – sablon	208
A.3. Szinkronizációs mátrix (Synchronization Matrix) – egyszerűsített sablon	208
A.4. Indikátor-katalógus (MOP/MOE) – sablon	209
A.5. Handover/átadás ellenőrzőlista (Shift handover checklist)	209
Melléklet B – SSA (Strategic Assessment) részletes szerkezete és kitöltési útmutató	210
B.1. Executive summary és döntési kérdés	210
B.2. Mandátum, felhatalmazás, politikai célok és korlátok	210
B.3. Műveleti környezet (OE) – rendszerleírás	211
B.4. Szereplők, érdekeltségek és viszonyok	211
B.5. Ellenség modell: COA-k, sebezhetőségek és adaptáció	211
B.6. End state, success criteria, exit criteria	212
B.7. SSA minőségbiztosítás: gyors checklist	212
Melléklet C – MRO/COA csomag: sablonok, pontozás és dokumentált bizonyíték	213
C.1. COA leíró sablon (COA narrative template)	213
C.2. Pontozási fegyelem: súlyozás és bizonyíték	213
C.3. Wargame dokumentálás: minimum outputok	214
Melléklet D – CONOPS → OPLAN átmenet: termékek, annex-index és végrehajthatósági tesztek	215
D.1. Annex-index (példa struktúra)	215
D.2. Végrehajthatósági tesztek (feasibility tests)	215
Melléklet E – DST/DSM módszertan: döntési pontok, triggererek és CCIR összerendelése	217
E.1. Döntési pont katalógus – „kevesebb, de jobb” elv	217
Melléklet F – Assessment mélyítés: adatforrások, confidence level és trianguláció	218

F.1. Adatforrás-típusok és tipikus torzítások	218
F.2. Confidence level skála – javasolt egyszerűsítés	218
F.3. Trianguláció a DC mérésében.....	218
Melléklet G – Gyors referenciakártyák (Phase 1–6) – tervezési minimum.....	219
G.1. Phase 1 (ISA) – minimum.....	219
G.2. Phase 2 (SSA) – minimum	219
G.3. Phase 3 (MRO/COA) – minimum	219
G.4. Phase 4 (CONOPS/OPLAN) – minimum	219
G.5. Phase 5 (Execution) – minimum	219
G.6. Phase 6 (Transition) – minimum.....	220
Melléklet H – Cyber és információs dimenzió integrációja a COPD-tervezésben	221
H.1. Cyber hatások és decisive conditions – mérhetőség és proxyk	221
H.2. Információs környezet és STRATCOM – LoE owner és indikátor-fegyelem	221
H.3. Jog és politika: acceptability „korai beégetése”	221
H.4. CCIR cyber/IO kontextusban – „mit nem tudunk, ami nélkül nem dönthetünk”	222
Melléklet I – Szemléltető tervezési vignetta: a COPD-termékek összekapcsolása egy rövid forgatókönyvben.....	223
I.1. Kiinduló helyzet (ISA/SSA vázlat)	223
I.2. Design vázlat: COG, DC-k, LoE-k.....	223
I.3. COA-k röviden és döntési logika	223
I.4. DST/DSM példa: reputációs DP és LOG DP	224
I.5. Assessment vázlat és tranzíció.....	224
RÖVIDÍTÉSEK JEGYZÉKE	225

1. BEVEZETÉS

A 21. századi biztonsági környezetet a többdimenziós (katonai–politikai–gazdasági–információs) versengés, a hibrid fenyegetések, valamint a technológiai forradalom felgyorsult üteme egyszerre alakítja. A NATO 2022-es Stratégiai Konceptiója kiemeli, hogy a Szövetségnek meg kell őriznie technológiai előnyét, és az „emerging and disruptive technologies” (EDT) – köztük a mesterséges intelligencia (MI) – egyszerre jelentenek lehetőséget és kockázatot [1]. Az Európai Unió Stratégiai Iránytűje (Strategic Compass) hasonlóan hangsúlyozza a hibrid fenyegetések, a kiber- és információs műveletek, valamint a technológiai fejlődésből fakadó kihívások kezelésének szükségességét [2]. Magyarország Nemzeti Biztonsági Stratégiája a „technológiai forradalom társadalomformáló hatásait” a változékony biztonsági környezet meghatározó tényezői közé sorolja, és a biztonságot kifejezetten kiterjeszti a technológiai és kibertérbeli dimenziókra is [3]. E stratégiai keretek együtt azt jelzik, hogy a döntéshozatali támogató, adat- és tudásintenzív képességek fejlesztése – ideértve az MI célzott alkalmazását – ma már nem opcionális „innováció”, hanem versenyképességi és műveleti hatékonysági követelmény.

A disszertáció témája: a mesterséges intelligencia alkalmazási lehetőségei a katonai stratégiai–hadműveleti szintű művelettervezés támogatásában, a NATO Comprehensive Operations Planning Directive (COPD) módszertanára építve. A COPD – amelyet az ACO (Allied Command Operations) tervezői közössége fejlesztett – a tervezők munkáját segítő útmutató; saját megfogalmazása szerint „a way, not the way”, és más NATO tervezési kiadványokat kiegészítve alkalmazandó [4]. A COPD evolúcióját áttekintő hazai szakirodalom ugyancsak rögzíti, hogy a dokumentum nem doktrína, hanem az ACO által kidolgozott tervezési útmutató, amely a változó környezethez igazodva folyamatosan fejlődött [5]. A kutatás alapfeltevése, hogy a COPD módszertani logikája alkalmas „keretet” ad az MI-támogatás rendszerezett azonosításához (mely tervezési tevékenységeknél, milyen MI-képességek, milyen adatokkal és kontrollokkal alkalmazhatók).

A NATO digitális transzformációs törekvései és az adatközpontú működés előtérbe kerülése közvetlenül érinti a tervezést: a NATO Digitális Transzformációs Megvalósítási Stratégiája a többdimenziós műveletek (multi-domain operations), az interoperabilitás, az adatmegosztás és az MI/adat felelős használatának fontosságát hangsúlyozza [6], miközben a NATO Adatstratégiája kifejezetten a „data-driven decision-making” és az egységesített

adatmegosztási ökoszisztéma megteremtését tekinti kulscélként [7]. Ezek a keretek a disszertáció számára kettős kiindulópontot jelentenek: egyrészt megerősítik a téma relevanciáját, másrészt olyan követelményi környezetet teremtenek (interoperabilitás, adat-kormányzás, felelős MI), amelyet az MI-alapú tervezéstámogatás nem kerülhet meg.

1.1. A kutatás időszerűsége és relevanciája (NATO/EU/MH kontextus)

A művelettervezés időszerűsége nemcsak a fenyegetések változékonyságából, hanem az információs és döntési ciklusok felgyorsulásából is fakad. A stratégiai–hadműveleti szinten a parancsnok és törzs feladata, hogy a politikai-stratégiai célokat műveleti célrendszerre, majd koherens műveleti elgondolássá alakítsa; e tevékenység egyszerre épít „művészetre” (operational art) és szisztematikus tervezési eljárásokra. Az AJP-5 kiemeli, hogy az operations design iteratív megértési és problémakeretezési folyamat, amely a parancsnok és törzs számára a működési környezet megértését és életképes műveleti megközelítések kialakítását támogatja [8]. A hasonló logika a JP 5-0-ban is megjelenik: a tervezés a parancsnok és törzs közötti iteratív párbeszédre épül, és a tervezési funkciók (stratégiai iránymutatás, koncepciófejlesztés, tervfejlesztés, tervértékelés) gyakran átfedésben, „párhuzamosan” futnak [9]. A tervezési komplexitás tehát nem csupán adminisztratív, hanem kognitív és szervezeti kihívás is.

A modern műveleti környezetben a döntések megalapozását – különösen a hibrid, információs és kibertérbeli dimenziókban – egyre inkább nagy mennyiségű, heterogén (strukturált/strukturálatlan) adat és szöveges információ határozza meg. A JIPOE (Joint Intelligence Preparation of the Operational Environment) doktrína definíciója szerint ez egy analitikai folyamat, amely hírszerzési értékeléseket és becsléseket állít elő a parancsnoki döntéshozatal támogatására, négy fő lépésben (OE meghatározása, hatások leírása, szereplők értékelése, valószínű ellenséges COA meghatározása) [10]. Mivel a COPD-tervezésben a helyzetértés és az ellenség/egyéb szereplők szándékainak előrejelzése meghatározó, az MI alkalmazhatóságának egyik fő terepe éppen a JIPOE jellegű elemzések támogatása (adatfúzió, trend- és mintázatelemzés, strukturálatlan szövegek feldolgozása).

Magyarország Nemzeti Katonai Stratégiája hangsúlyozza, hogy a megújuló Magyar Honvédség a NATO kollektív védelem keretében, ugyanakkor erős nemzeti önerőre támaszkodva garantálja az ország biztonságát és a szövetségi hozzájárulások hitelességét [11]. Ebből a nézőpontból a NATO-kompatibilis tervezési képességek fejlesztése – beleértve az MI

alapú döntéstámogatás integrálását – egyszerre szolgálhat nemzeti és szövetségi célokat: (i) gyorsabb és konzisztens(ebb) tervtermék-előállítást; (ii) szélesebb helyzetértést a komplex OE-ben; (iii) interoperábilis, adatmegosztásra épülő munkafolyamatokat; (iv) a felelős, auditálható MI-használatból fakadó bizalomerősítést.

1.2. Tudományos problémafelvetés és kutatási rés (research gap)

A COPD-alapú művelettervezés a gyakorlatban jelentős mennyiségű információgyűjtést, elemzést, szintézist és formalizált tervdokumentációt igényel. Miközben a civil szférában az MI-alapú döntéstámogatás és szövegfeldolgozás tömegessé vált, a katonai tervezésben az MI alkalmazása gyakran szigetszerű, eseti jellegű (pl. adatelemzés egy-egy részfeladatra), és ritkán illeszkedik szabatosan a tervezési módszertan (jelen értekezésben: COPD) tevékenységeihez és termékeihez. A tudományos probléma ezért így fogalmazható meg:

Hogyan lehet az MI-képességeket úgy azonosítani, értékelni és integrálni, hogy azok a COPD szerinti tervezési folyamat meghatározott lépéseiben mérhetően javítsák a tervezés hatékonyságát és minőségét, miközben a katonai alkalmazás kockázatai (félrevezetés, torzítás, adatbiztonság, jogi/etikai megfelelés) kontrollálhatók maradnak?

A kutatási rés (research gap) három, egymással összefüggő hiányosságra épül:

1. Módszertani illesztés hiánya: kevés olyan átfogó keret található, amely a COPD (és kapcsolódó NATO tervezési logika) konkrét tevékenységeit, döntési pontjait és tervtermékeit MI-támogatási lehetőségekkel párosítja [4], [5].
2. Értékelési és mérési hiány: a katonai tervezésben ritkán jelennek meg egységes, empirikusan tesztelhető indikátorok arra, hogy az MI milyen mértékben gyorsítja a tervezést, csökkenti-e a kognitív terhelést, illetve hogyan hat a termékminőségre (koherencia, konzisztencia, visszakövethetőség).
3. Felelős és megbízható alkalmazás operacionalizálatlansága: a „trustworthy / responsible AI” elvek megjelennek a nemzetközi kockázatkezelési keretekben, de a tervezési folyamatba való beépítésük gyakorlati módszertana nem kiforrott. E tekintetben a NIST AI RMF a kockázatkezelés (GOVERN–MAP–MEASURE–MANAGE) struktúrájával és a „trustworthiness” attribútumokkal hasznos hivatkozási alapot ad [12], míg a DoD felelős MI

útvonalterve a katonai környezetben elvárt bizalom- és kormányzási dimenziókat konkretizálja [13].

1.3. Kutatási kérdések

A disszertáció a fenti problémafelvetést az alábbi kutatási kérdésekre bontja:

K1. A COPD tervezési folyamata (tevékenységek, döntési pontok, tervtermékek) közül mely elemek alkalmasak MI-támogatásra, és milyen feltételek mellett? [4]

K2. Milyen MI-módszerek (pl. NLP, tudásgráfok, gépi tanulós előrejelzés, szimuláció és generatív modellek) illeszthetők a tervezés különböző feladataihoz (pl. helyzetértés, COA-fejlesztés, kockázatelemzés, dokumentálás), és milyen adat- és integrációs követelményekkel? [6], [7]

K3. Milyen kockázatok (torzítás, megtévesztés, „hallucináció”, információbiztonság, jogi/etikai megfelelés) jelennek meg a COPD-alapú MI-támogatás során, és ezek hogyan kezelhetők egy felelős AI-kormányzási keretben? [12], [13]

K4. Milyen mérőszámokkal és értékelési eljárásokkal igazolható, hogy az MI-támogatás javítja a tervezési ciklus hatékonyságát és a tervtermékek minőségét (különösen NATO/MH alkalmazási környezetben)?

1.4. Hipotézisek

A kutatás a következő – empirikusan vizsgálható – hipotéziseket fogalmazza meg:

H1. A COPD dokumentum-központú tervtermékeinek előállításában az MI-alapú szövegfeldolgozás és tudásmenedzsment (pl. automatikus kivonatolás, entitás- és relációkinyerés, hivatkozás-konzisztencia ellenőrzés) mérhetően csökkenti a tervezés átfutási idejét, miközben a termékek strukturális megfelelését javítja.

H2. A hírszerzési megalapozást (különösen JIPOE jellegű elemzéseket) támogató adatfúziós és mintázatfelismerő MI-megoldások növelik a helyzetértés mélységét és csökkentik a kognitív torzításokból eredő hibák valószínűségét, feltéve hogy a folyamat humán „human-in-the-loop” ellenőrzéssel működik [10].

H3. Megfelelő kormányzás és ellenőrzés hiányában az MI-támogatás növeli a megtévesztés, a hibás következtetés és a nem kívánt információszivárgás kockázatát; e kockázatokat csak kockázat-alapú (GOVERN–MAP–MEASURE–MANAGE) keret, auditálhatóság és szigorú adatkezelés mellett lehet elfogadható szintre csökkenteni [12], [13].

H4. A COPD-hez illesztett, moduláris, interoperábilis MI-támogató architektúra – amely a NATO digitális transzformációs és adatstratégiai irányokhoz igazodik – növeli a bevezethetőség és elfogadottság esélyét a NATO/EU/MH kontextusban [6], [7], [11].

1.5. Kutatási célok és célkitűzések

A disszertáció általános célja olyan tudományos és gyakorlati keret kidolgozása, amely a COPD szerinti művelettervezésben az MI alkalmazhatóságát rendszerezett módon feltárja, és javaslatot ad a bevezetés feltételeire és kockázatkezelésére. Ennek keretében a fő célkitűzések:

1. Fogalmi és folyamatleíró cél: a COPD tervezési logika (tevékenységek, döntések, termékek) olyan bontása, amelyhez MI-támogatási „belépési pontok” rendelhetők [4], [5].
2. Alkalmazási katalógus: MI-use case-ek és képességek katalógusának létrehozása COPD-fázisok szerint (pl. információ-előkészítés, COA-fejlesztés, tervszövegezés).
3. Felelős MI-követelmények: a katonai MI-támogatás kockázatainak és kontrolljainak levezetése, a nemzetközi kockázatkezelési és felelős MI keretek adaptálásával [12], [13].
4. Értékelési keret: mérőszám- és módszertan-javaslat az MI-támogatás hatásának értékelésére (idő, minőség, konzisztencia, visszakövethetőség, bizalom).
5. MH relevancia: ajánlások megfogalmazása a Magyar Honvédség NATO-kompatibilis művelettervezési és digitális fejlesztési törekvéseihez illeszkedően [11], [3].

1.6. Alkalmazott módszerek és adatforrások

A kutatás több módszer kombinációjára épül (mixed-methods), mert a művelettervezési folyamat egyszerre szabályozott (dokumentumok, eljárások) és szocio-technikai (emberi döntéshozatal, szervezeti kultúra) jelenség.

Irodalom- és dokumentumelemzés: NATO tervezési alapidokumentumok (COPD, AJP-5), kapcsolódó amerikai tervezési/hírszerzési doktrínák (JP 5-0, JP 2-01.3) és a releváns stratégiai keretek (NATO SC, EU Compass, NATO DTIS, NATO Data Strategy, magyar NBS/NKS) elemzése [4], [1]–[3], [6]–[7], [11], [8]–[10].

Komparatív elemzés: a COPD és más tervezési keretek (AJP-5; JP 5-0) koncepcionális összevetése az MI-támogatási pontok szempontjából [8], [9].

Esettanulmány (scenario-based case study): egy (nyilvános, nem minősített) válságforgatókönyvön keresztül a COPD-lépésekhez rendelt MI-támogatási use case-ek „végigjátszása”, a szükséges adatok és kontrollok azonosításával.

Szakértői interjúk / Delphi-elemek: tervezői, hírszerzési, informatikai és jogi/etikai szakértők bevonása az alkalmazhatóság, kockázatok és szervezeti bevezethetőség validálására.

Tervezett (korlátozott) prototípus / demonstrátor: ahol lehetséges, kis kockázatú (pl. szöveg-összefoglalás, követelmény-kivonatolás) MI-funkciók kipróbálása nyílt adatokon, a mérőszámok teszteléséhez.

Adatforrások: (i) hivatalos NATO/EU/magyar stratégiai és doktrinális dokumentumok; (ii) nyílt tudományos publikációk (MI a C2-ben, döntéstámogatásban, hadműveleti elemzésben); (iii) nyílt esettanulmányok és szimulációs/forgatókönyv-anyagok; (iv) szakértői interjúk anonimizált jegyzőkönyvei.

1.7. Lehatárolások, feltételezések, korlátok

A disszertáció fókusza a stratégiai–hadműveleti tervezési szint (NATO COPD) MI-támogatása. Ennek megfelelően:

Lehatárolás: nem cél a taktikai szintű harcászati döntéstámogató rendszerek vagy fegyverrendszerek MI-jének elemzése; a vizsgálat a törzsszintű tervezésre és a tervtermékekre koncentrál.

Minősítési korlát: a kutatás nyílt forrásokra és nem minősített esettanulmányokra támaszkodik; a minősített rendszerek és adatok részletei csak általánosítható módon (követelmények szintjén) jelennek meg.

Technológiai feltételezés: a vizsgált MI-megoldásoknál feltételezett a humán ellenőrzés (human-in-the-loop), valamint az auditálhatóság és adatkezelési megfelelés alapkövetelménye [12], [13].

Korlát: a generatív modellek esetén a hibás tartalom („hallucináció”) és a forrásmegjelölés bizonytalansága külön kontrollokat igényel; ezért a disszertáció a kritikus döntési pontoknál a döntéshozói felelősség és ellenőrzés elsődlegességét tekinti irányadónak.

1.8. Az értekezés felépítése és logikája

Az értekezés logikája a (1) probléma–környezet, (2) COPD-folyamatelemzés, (3) MI-képességek és use case-ek, (4) kockázatkezelés és kormányzás, (5) értékelés és ajánlások ívét követi.

2. TAXONÓMIAI KERET – A MESTERSÉGES INTELLIGENCIA ÉS FELTÖREKVŐ/DISZRUPTÍV TECHNOLÓGIÁK A KATONAI KÖRNYEZETBEN

Az elméleti keret célja, hogy egységes fogalmi és rendszertani alapot biztosítson a mesterséges intelligencia (MI) katonai alkalmazásainak – különösen a stratégiai–műveleti tervezési folyamatot támogató megoldásoknak – vizsgálatához. A fejezet először a NATO által használt „feltörekvő és diszruptív technológiák” (Emerging and Disruptive Technologies, EDT) koncepcióját és az ahhoz kapcsolódó intézményi/stratégiai kereteket tekinti át, majd az MI alapfogalmait és taxonómiáját rendezi, végül katonai funkcionális területek szerint mutat be alkalmazási mintázatokat. A későbbi alfejezetek (2.4–2.6) erre építve tárgyalják a megbízhatósági, magyarázhatósági, adatminőségi, kiberbiztonsági és emberi kontrollkövetelményeket.

2.1. Feltörekvő és diszruptív technológiák (EDT) – fogalom és NATO-értelmezés

Az „emerging and disruptive technologies” (EDT) fogalompár a biztonság- és védelempolitikai szakirodalomban a 2010-es évek végétől vált meghatározóvá. A kifejezés a technológiai fejlődés gyorsulására és a technológiák konvergenciájára adott stratégiai válasz: a cél nem pusztán az újdonságok felsorolása, hanem annak megértése, hogy a technológiai trendek miként alakítják át a katonai képességeket, a hadviselés karakterét, valamint a politikai–katonai döntéshozatal és a műveleti tervezés mozgásterét.

A NATO a 2020–2040 időhorizontot vizsgáló S&T Trends jelentésében (STO) az „emerging” technológiákat szűk értelemben olyan fejlesztésekként/ tudományos felfedezésekként definiálja, amelyek várhatóan a 2020–2040 időszakban érnek el érettséget, és jelenleg még nem elterjedtek, illetve hatásuk a szövetségi védelemre és elrettentésre nem teljes mértékben kiaknázott [2, p. 12]. A „disruptive” jelző ehhez képest a katonai hatás oldaláról ragadja meg a jelenséget: olyan technológiai ugrást jelöl, amely képes a meglévő képesség–ellentevékenységi egyensúlyokat felborítani, új sebezhetőségeket létrehozni, és ezzel az erőalkalmazás módját, költségfüggvényeit vagy időigényét alapvetően megváltoztatni.

A NATO témalapú összefoglalója szerint az EDT-k egyszerre jelentenek kockázatot és lehetőséget: egyrészt gyorsabb, rugalmasabb és költséghatékonyabb katonai képességeket tehetnek elérhetővé, másrészt az ellenfelek is kihasználhatják őket, illetve a civil társadalom ellen is alkalmazhatók (hibrid fenyegetések, információs műveletek, kritikus infrastruktúra elleni támadások) [14]. A NATO ezért kettős logikát követ: „foster and protect” – azaz egyrészt ösztönözni kívánja a duál felhasználású innovációkat és azok gyors adaptációját, másrészt védeni akarja saját innovációs ökoszisztémáját az ellenséges befolyásolástól, manipulációtól és technológiai kiszivárgástól [14].

2.1.1. A NATO EDT-konceptió intézményi beágyazottsága és politikai–katonai célrendszere

A NATO 2019–2025 közötti időszakban az EDT-k köré épülő intézményi és szabályozási architektúrát fokozatosan bővítette. A NATO témalapú áttekintése szerint a Szövetség kilenc kiemelt EDT-területre fókuszál: mesterséges intelligencia, autonóm rendszerek, kvantumtechnológiák, biotechnológia és emberi képességfokozás, űr, hiperszonikus rendszerek, új anyagok és gyártástechnológiák, energia és hajtóművek, valamint a következő generációs kommunikációs hálózatok [14]. Az EDT-k kezelése tehát nem egyetlen szűk szakpolitikai kérdés, hanem a képességfejlesztés, a védelmi ipar, az interoperabilitás és a digitális transzformáció metszéspontjában áll.

A „gyors adaptáció” (rapid adoption) a NATO 2025-ös összefoglalója szerint explicit cél: a Szövetség vezetői a 2025-ös hágai csúcson olyan akciótervet fogadtak el, amely a technológiai termékek szövetségi adaptációját „maximum 24 hónapon belül” kívánja megvalósítani [14]. A gyorsítás logikája közvetlenül érinti a katonai műveleti tervezés támogatását is: a tervezési eljárások (pl. COPD) és a támogató informatikai rendszerek (C2/decision support) csak akkor maradnak relevánsak, ha képesek integrálni az új, duál felhasználású képességeket – miközben a biztonsági akkreditáció, a minősített adatkezelés és az interoperabilitási standardok betartása nem sérül.

A NATO az EDT-k gyorsítására több eszközt hozott létre: a DIANA (Defence Innovation Accelerator for the North Atlantic) innovációs gyorsítót és a NATO Innovation Fundot (többszuverén kockázatitőke-alap) a transzatlanti innovációs ökoszisztéma erősítésére pozicionálva [14]. A kifejezetten MI-hez kapcsolódó intézményi elem a NATO Data and Artificial Intelligence Review Board (DARB), amely a NATO AI stratégia által rögzített felelős MI-

használati elvek (Principles of Responsible Use) operacionalizálását szolgálja, és többek között egy „Responsible AI certification standard” fejlesztését is napirendre tűzte [14], [15].

E fejlemények stratégiai üzenete a műveleti tervezés szempontjából kétirányú. (1) A döntésciklusok gyorsulása (adat→információ→döntés→hatás) egyre inkább a technológiai előny függvénye. (2) A technológiai előny nem csak a harcászati eszközökben jelenik meg, hanem a törzsmunkát támogató információs rendszerekben is: a COA-k (Courses of Action) generálása, a kockázatértékelés, a logisztikai „what-if” elemzések, valamint a hatásalapú tervezés (effects-based thinking) mind olyan területek, ahol a nagy adathalmazok és a tanuló algoritmusok képesek csökkenteni az elemzői terhelést és növelni a konzisztenciát – feltéve, hogy a szakmai felelősség és az emberi kontroll megfelelően érvényesül.

NATO kilenc kiemelt feltörekvő és diszruptív technológiai területe (EDT)

Mesterséges intelligencia (AI)	Autonóm rendszerek	Kvantumtechnológiák
Biotechnológia és emberi képességfokozás	Űr (Space)	Hiperszonikus rendszerek
Új anyagok és gyártástechnológiák	Energia és hajtóművek	Következő generációs kommunikációs hálózatok

2-1. ábra – NATO kilenc kiemelt EDT-területe (saját szerkesztés a NATO összefoglalója alapján [14]).

2.1.2. EDT-k és a műveleti tervezés: releváns hatásmechanizmusok

Az EDT-k hatása a műveleti tervezésre nem kizárólag egy-egy technológia „bevezetésében” ragadható meg, hanem a tervezési környezet strukturális változásában. A NATO STO jelentés hangsúlyozza, hogy az EDT-k azonosításának egyik szempontja, hogy a fejlesztés „jelentősen

befolyásolja a szövetségi képesség- vagy tervezési döntéseket” [2, p. 9]. Ez különösen fontos COPD-alapú tervezésben, ahol a problémakeretezés (mission analysis), a COA-fejlesztés és -értékelés, valamint a kockázat- és erőforrás-analízis idő- és adatintenzív folyamat.

Az EDT-k által kiváltott tervezési kihívások közül a következőket célszerű kiemelni:

- Információs túltelítettség és érzékelési forradalom: a többdoménű (MDO) hadviselésben a szenzorok és nyílt források (OSINT) által termelt adatmennyiség exponenciálisan növekszik; a szűrés és megbízhatósági értékelés egyre inkább automatizálásra szorul [15, pp. 30–33].
- Döntési időablakok szűkülése: a hiperszonikus rendszerek, a nagy sebességű kibertámadások és a spektrumbeli műveletek következtében a reakcióidők rövidülnek, ami a „döntési fölény” (decision advantage) elérését a C2-rendszerek és a döntéstámogatás szintjére emeli [2, pp. 1–6].
- Rendszerkonvergencia és összekapcsoltság: az EDT-k nem izoláltan jelennek meg, hanem egymást erősítő módon (pl. MI + autonómia + új kommunikációs hálózatok + úralapú ISR) – ami a tervezésben integrált, többdoménű hatásláncok (effects chains) modellezését igényli [2, pp.10–11].
- Interoperabilitás és standardizáció dilemmája: a gyors adaptáció és a kísérletezés a NATO-ban csak akkor skalázható, ha a technológiai megoldások interoperábilisak; ugyanakkor a túl korai standardizáció gátolhatja az innovációt. E feszültség kezelésére jelennek meg olyan intézményi mechanizmusok, mint a DIANA tesztközpont-hálózata és a DARPA felelős MI-tanúsítási törekvései[14],[15].

A katonai tervezés szempontjából az EDT-k „operacionalizálásának” kulcsa a képességek feladat–funkció–hatás leképezése (capability-to-task mapping). A NATO ACT MUAAR kezdeményezés (Military Uses of Artificial Intelligence, Automation, and Robotics) kifejezetten azt a célt fogalmazza meg, hogy egy kompakt feladatlistát (task compendium) adjon arról, mely katonai feladatok profitálhatnak az AI/automatizálás/robotika kombinációjából, és mely területeken a legnagyobb a „payback” hatékonyság, költségmegtakarítás vagy letalítás szempontjából [7, p. 1]. A tervezési folyamatok akkor tudják ezt a logikát kiaknázni, ha a törzsek rendelkeznek (1) a releváns adatforrásokhoz való hozzáféréssel, (2) megfelelően akkreditált és auditálható algoritmusokkal, valamint (3) olyan eljárásokkal, amelyek az MI-ajánlásokat az emberi döntéshozatalba beágyazzák.

2.2. MI alapfogalmak és taxonómia (ML/DL, generatív MI, hibrid rendszerek)

Az MI-nek annak ellenére, hogy már az ötvenes évek óta használt fogalom, ma sincs általánosan elfogadott tudományos definíciója. Megállapításom az, hogy az MI nem értelmezhető önálló alkalmazásként, hanem olyan technológia, amely meglévő funkcionális megoldásokat támogat meghatározott problémák megoldására kifejlesztett algoritmusokon alapuló eljárásokkal. Ezen algoritmusok nagy adatkészletek összegyűjtésére, rendszerezésére, feldolgozására, elemzésére, továbbítására és az ezekre való reagálásra alkalmasak, azaz képesek az emberi értelem kognitív képességének megfelelő, illetve azt közelítő műveletekre, mégpedig nagyobb sebességgel.

Az MI-nek alapvetően három típusát különböztetjük meg:

- Narrow (AI) MI, azaz szűk vagy gyenge mesterséges intelligencia, amely olyan számítógépes rendszer, amely az embernél hatékonyabban el tud végezni egy pontosan meghatározott feladatot. Itt tartunk ma.
- General (AI) MI, általános mesterséges intelligencia, amelyet olykor „erős MI-nek” is neveznek, az ember kognitív képességeit meghaladva képes bármilyen intellektuális feladatot elvégezni. Az ilyen típusú MI-vel működő robotokat láthatunk filmekben, ahol tudatos gondolkodással saját céljaiknak megfelelően működnek. Ez ma még nagyrészt a fantázia világa.
- Artificial Super Intelligence – ASI, a mesterséges szuperintelligenciával rendelkező számítógép képes az embert minden területen felülmúlni, például akár a tudományos kutatásban, általános bölcsességben és a társadalmi képességekben is. Erről a tudósok jó része meg van győződve, hogy elérhetetlen.

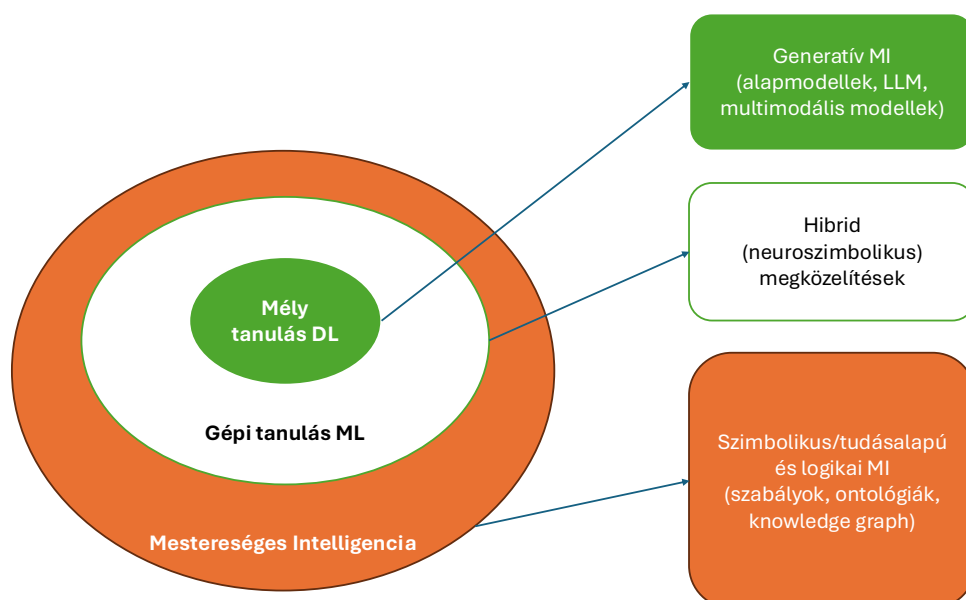
Az MI fogalma a civil és katonai diskurzusban egyszerre tudományos-technikai és stratégiai-politikai kategória. A disszertáció szempontjából olyan definíció szükséges, amely egyrészt kompatibilis a nemzetközi szervezetek szabályozási nyelvével, másrészt elég részletes ahhoz, hogy a COPD-alapú tervezést támogató konkrét alkalmazások (pl. COA-generálás, kockázatmodellezés, OSINT feldolgozás) elhelyezhetők legyenek.

Az Európai Bizottság AI Act-hez kapcsolódó melléklete (Annex I) az MI-technikai megközelítéseket három nagy csoportba sorolja: (A) gépi tanulás (felügyelt, felügyelet nélküli, megerősítéses tanulás; beleértve a mélytanulást), (B) logikai- és tudásalapú megközelítések (szabályok, tudásreprezentáció, deduktív/inferenciamechanizmusok, szimbolikus érvelés, szakértői rendszerek), valamint (C) statisztikai megközelítések (Bayes-bebecslés, keresés és

optimalizálás) [12, p. 2]. Ez a csoportosítás – bár eredetileg szabályozási célú – jól használható katonai taxonómiaként is, mert világosan különválasztja a „tanuló” (data-driven) és a „tudásalapú” (rule/knowledge-driven) rendszereket, illetve helyet ad a hibrid (neuro-szimbolikus) megoldásoknak.

A NATO 2024-ben frissített AI stratégiájának összefoglalója külön nevesíti a generatív MI-t és az MI-alapú információs eszközöket mint új fejleményeket, amelyekhez a szövetségi elvek és alkalmazási keretek igazítása szükséges [16]. A generatív modellek (különösen a nagy nyelvi modellek – LLM-ek) megjelenése azért lényeges katonai kontextusban, mert (1) a természetes nyelvű információk feldolgozását és a törzskommunikációt is érintik, valamint (2) a döntéstámogatásban új felhasználási mintázatokat nyitnak (pl. össz-adatforrású összefoglalás, strukturált törzsdokumentumok előállítás, alternatív COA-vázlatok generálása). Ugyanakkor e modellek kockázatai (hallucináció, adatszivárgás, prompt-injekció) a katonai alkalmazhatóságot szigorú kontrollokhoz kötik – ezek részletes elemzése a 2.4–2.5 alfejezetekben történik.

A következőkben az MI taxonómiáját három, egymást kiegészítő tengely mentén rendezzük: (1) tanulási paradigma (felügyelt/felügyelet nélküli/megerősítéses), (2) modellarchitektúra (klasszikus ML vs. DL vs. alapmodellek), (3) tudásreprezentáció és érvelés (szimbolikus vs. hibrid).



2-2. ábra – MI taxonómia (saját szerkesztés; EU AI Act Annex I kategóriái és NATO AI-stratégia alapján [17], [16]).

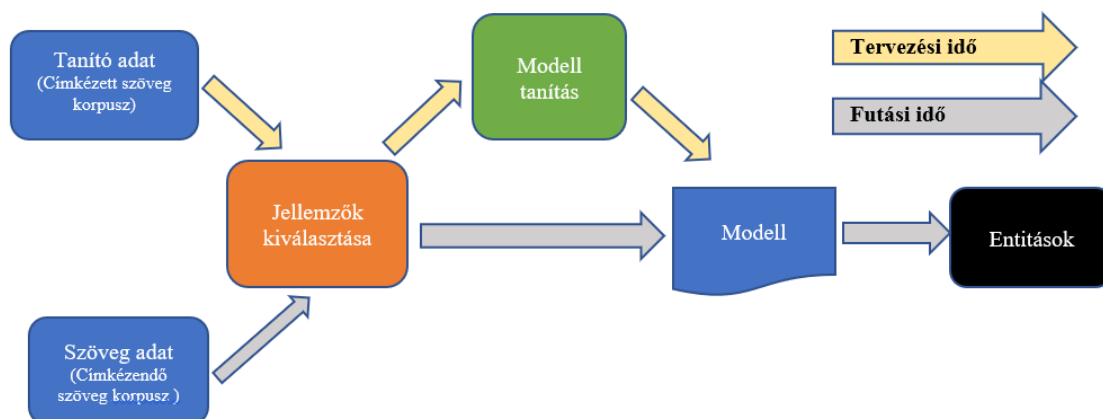
2.2.1. Gépi tanulás (ML): felügyelt, felügyelet nélküli és megerősítéses tanulás

Az MI egy részhalmazának tekinthető a gépi tanulás (*Machine Learning*, ML) amely matematikai adatmodellekkel tanít be számítógépeket közvetlen felügyelettel vagy anélkül. A gépi tanulás algoritmusokkal azonosít mintákat az adatokban, amelyekből adatmodellt készít, majd előrejelzéseket és válaszokat ad.

A gépi tanulás olyan módszerek gyűjtőneve, amelyekben a rendszer a rendelkezésre álló adatokból becsüli meg a döntési szabályokat vagy a predikciós függvényeket. Katonai környezetben – különösen az ISR és a kiber területén – az ML általában „mintázatfelismerési” és „anomáliadetektálási” feladatokban jelenik meg, ahol az emberi elemző kapacitása nem képes lépést tartani a szenzor- és hálózati adatok volumenével.

Felügyelt tanulás (supervised learning, SL) esetén a tanító adatok címkézettek (pl. „baráti/ellenséges emisszió”, „hamis/hiteles tartalom”, „típusazonosítás”). A katonai alkalmazások előnye, hogy a modell kimenete kontrollálható a címkék minősége révén; hátránya, hogy a minősített környezetben a címkézett adat előállítása költséges és szakértői időt igényel.

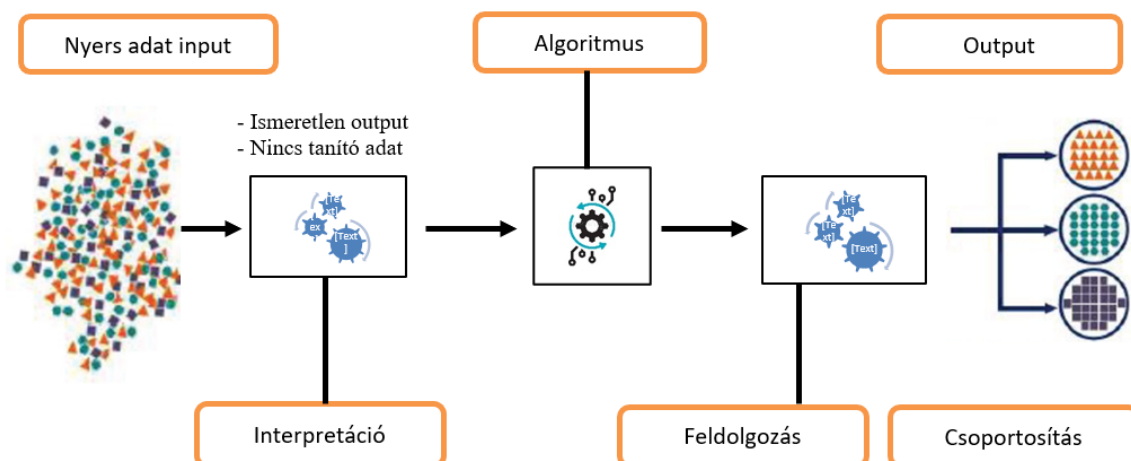
A felügyelt tanulás során az osztályozó paramétereit az ismert kategóriákból álló minták felhasználásával a kívánt teljesítmény eléréséhez igazítják. Az SL egy funkcionális gépi tanulási feladatot képez a címkézett tanító adatokból, amelyek tanító példákat tartalmaznak. Az SL-ben minden példa egy bemeneti objektumból (általában egy vektorból) és egy várható kimeneti értékből (felügyelt jelből) áll. Ezt mutatja be az alábbi, 2.3. ábra.



2.3. ábra: A felügyelt tanulás (SL) diagramja. Forrás: Wei Wang et al.: *Investigation on Works and Military Applications of Artificial Intelligence. IEEE Access*, 8. (2020), 131614–131625. alapján a szerző fordítása

Felügyelet nélküli tanulás során a tanító adatok nincsenek címkézve, a tanulási cél a megfigyelt értékek osztályozása vagy megkülönböztetése. Lényegében ez egy statisztikai módszer, amely képes felismerni a jelöletlen adatok potenciális struktúráit. A felügyelet nélküli tanulás diagramját a 2. ábra mutatja. Az UL-t gyakran használják az adatbányászatban, hogy feltárjanak valamit nagy mennyiségű, strukturálatlan adatban. Ilyen például a képfelismerés.

Felügyelet nélküli tanulás (unsupervised learning, UL) célja a struktúrák, klaszterek, rejtett dimenziók feltárása címkézés nélkül. OSINT és SIGINT feldolgozásban tipikus a témamodellezés, klaszterezés és a dimenziócsökkentés, amely az elemzők számára „jelölt” (candidate) anomáliákat vagy trendeket emel ki [15, pp.37–41].



2.4. ábra: A felügyelet nélküli tanulás (UL) diagramja, Forrás: Wang et al. (2020): i. m. alapján a szerző fordítása

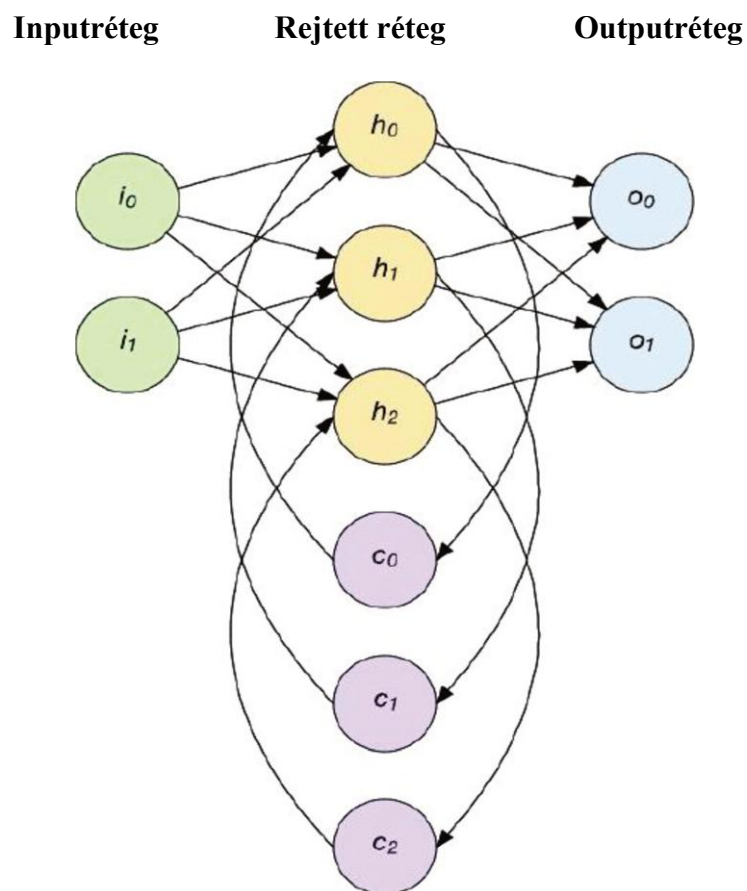
Az ML fejlettebb szintje a megerősítő gépi tanulás (reinforcement learning, RL), amikor a rendszert pozitív visszacsatolásokkal erősítik meg a felismerésekben. Az RL-t a kontrollelméletből (*control theory*), a statisztikából, a pszichológiából és kapcsolódó tárgyakból fejlesztették ki, és Pavlov feltételesreflex-kísérletére vezethető vissza.

A Megerősítő tanulás az akció–jutalom logikára épül: a rendszer környezettel interakcióban tanulja meg, milyen döntések maximalizálják a hosszú távú jutalmat. Katonai tervezési

támogatásban a RL elsősorban szimulációs környezetben (wargaming) és optimalizációs problémákban releváns (pl. erőforrás-allokáció, útvonaltervezés, ütemezés), ahol a „jutalom” a siker kritériumainak megfelelő, formalizált célfüggvény [2, pp. 9–10]. A RL gyakorlati kihívása, hogy a szimuláció minősége és a jutalomfüggvény helyes specifikációja meghatározza, a modell mennyire tanul „valós” katonai logikát.

2.2.2. Mélytanulás (DL) és alapmodellek: a generatív MI katonai jelentősége

A mélytanulás a neurális hálózatok többrétegű architektúráira épülő tanulási módszerek gyűjtőfogalma. A mély tanulás (Deep Learning, DL) esetében a gépet nagy mennyiségű adattal tanítják be összetett feladatokra az emberi agy analógiájára létrehozott neurális hálózatok segítségével, amelyben a neuronok (node-ok) egy-egy részfunkció végrehajtását végzik, illetve összegzik azokat. Itt fontos megemlíteni az úgynevezett Black Box jelenséget, amelynél az egyes neuronszintekben (hidden layer) végbemenő folyamatot az ember már nem képes követni, illetve átlátni, így azok jelentős megbízhatósági kockázatot hordoznak magukban.



Context nódok

2.5.. ábra: Egy neurális háló strukturális diagramja. Forrás: Wang et al. (2020): i. m.

A DL különösen ott vált dominánssá, ahol nagy dimenziójú, strukturálatlan adatok (kép, videó, hang, szöveg) feldolgozása szükséges – ez katonai kontextusban tipikusan ISR (kép- és videóanalitika), információs műveletek (tartalomosztályozás) és kiber (log-anomália) területeket érint.

A 2024-ben frissített NATO AI-stratégia összefoglalója kifejezetten utal arra, hogy az AI-ökoszisztéma gyors fejlődése – különösen a generatív MI és az AI-alapú információs eszközök – indokolja a stratégiai keretek aktualizálását [16]. A generatív MI (Generative AI) alatt a valószínűségi-neurális modelleket értjük, amelyek új tartalmakat (szöveg, kép, hang, kód) állítanak elő a tanulási eloszlás alapján. A katonai alkalmazás szempontjából fontos megkülönböztetés:

- „Alapmodellek” (foundation models): nagy, általános célú modellek, amelyeket később finomhangolnak (fine-tuning) vagy utasításkövetésre (instruction) állítanak.
- Nagy nyelvi modellek (LLM): természetes nyelvű következtetés, összefoglalás, információkinyerés.
- Multimodális modellek: több adatmodalitás (szöveg+kép+hang) integrált kezelése.

Tervezéstámogatásban a generatív modellek legkézenfekvőbb haszna a törzsmunkában jelentkező „dokumentum-intenzitás” csökkentése: helyzetértékelések, hírszerzési összefoglalók, COA-vázlatok, kockázati listák, valamint döntés-előkészítő briefek előállításának gyorsítása. A korlátok ugyanakkor katonai környezetben erősen érvényesülnek: a hallucináció (valótlan, de hihető állítás), a forrás-átláthatatlanság, a minősített adatkezelés és a biztonsági sebezhetőségek a generatív MI-t csak kontrollált, auditálható környezetben teszik alkalmazhatóvá [16],[10,pp.12–17].

A DL és alapmodellek katonai értelmezéséhez hasznos a „modell-életciklus” szemlélet: a tanítás (training), a validálás, a telepítés (deployment) és az üzemeltetés (monitoring) közötti váltások jelzik, hol vannak a kritikus pontok (adatminőség, drift, adversarial manipuláció). A NATO gyors adaptációs törekvései [14] e pontokon akkor találkoznak a valósággal, amikor a kísérleti prototípusból minősített, interoperábilis és kiberbiztos képességet kell létrehozni.

2.2.3. Szimbolikus és hibrid (neuro-szimbolikus) rendszerek: miért relevánsak a tervezésben?

A katonai döntés- és tervezéstámogatás klasszikus területei (szakértői rendszerek, szabályalapú érvelés, optimalizáció) sok esetben jobban illeszkednek szimbolikus vagy hibrid MI-megközelítésekhez, mint a tisztán adatalapú, „feketedoboz” jellegű DL-modellekhez. Az EU AI Act Annex I is külön kategóriaként kezeli a logikai- és tudásalapú megoldásokat, valamint a statisztikai/optimalizációs módszereket [12, p. 2].

A hibrid rendszerek lényege, hogy a tanuló komponensek (pl. NLP-kinyerés, mintázatfelismerés) eredményeit szimbolikus reprezentációkba (ontológiák, tudásgráfok, szabályok) emelik át, majd ezen a rétegen magyarázhatóbb érvelést vagy következtetést végeznek. Stratégiai–műveleti tervezésben ez azért fontos, mert a tervezési logika „szabályok és korlátok” mentén működik: erőképesség-korlátok, ROE, logisztikai kapacitás, politikai korlátozások, jogi megfelelés stb. A tisztán generatív modellek hajlamosak e korlátokat elnagyolni, míg a szimbolikus/hibrid keretrendszerben a korlátok explicit módon érvényesíthetők.

A Négyesi–Szabadszöldi szerzőpáros egy, a stratégiai OPLAN-t támogató MI-megoldásokat tárgyaló tanulmányban szintén hangsúlyozza, hogy a katonai döntéstámogató rendszerekben több, eltérő MI-paradigma alkalmazható: neurális hálók, Bayes-féle neurális hálók, fuzzy logika, genetikusan algoritmusok és szakértői rendszerek különböző feladatokra alkalmasak [16, p. 30]. Ez az „eszköztár-szemlélet” a disszertáció későbbi, COPD munkafázisokhoz illesztett MI-térképében (5. fejezet) közvetlenül hasznosítható.

2.3. Katonai alkalmazási területek (ISR, C2, EW, logisztika, cyber)

A katonai MI-alkalmazások rendszerezése többféle tengely mentén lehetséges: képességterületek (pl. ISR, C2), domének (land/air/maritime/space/cyber), vagy a tervezés–végrehajtás–értékelés ciklusa mentén. A disszertáció céljaihoz illeszkedően ebben az alfejezetben a NATO és a szakirodalom által gyakran használt funkcionális felosztást (ISR, C2, EW, logisztika, cyber) követjük, de minden területnél kiemeljük, milyen módon jelenhet meg a támogatás a stratégiai–műveleti tervezésben (COPD) – például helyzetértékelés, COA-fejlesztés, erőforrás- és kockázat-analízis, valamint a visszacsatolás (assessment) folyamatában.

A NATO ACT MUAAR kezdeményezés a katonai feladatok széles körére kívánja azonosítani, hol adhat hozzáadott értéket az AI/automatizálás/robotika „blendje”, külön kiemelve az elektromágneses spektrumot, az integrált lég- és rakétavédelmet, a logisztikát, az űrt és a kibertér, valamint a hagyományos doméneket (air/land/maritime) [7,p.1]. E kezdeményezéshez illeszkedően a következő részekben az alkalmazásokat „feladatcsomagokként” mutatjuk be: adatgyűjtés, feldolgozás–értelmezés, döntéstámogatás–optimalizáció, végrehajtás és reziliencia.

	Érzékelés & adatgyűjtés	Feldolgozás & értelmezés	Döntéstámogatás & optimalizálás	Végrehajtás & hatáskifejtés	Védelem & reziliencia
ISR/OSINT	✓	✓	✓		✓
C2 és tervezés		✓	✓	✓	✓
EW	✓	✓	✓	✓	✓
Logisztika és fenntartás	✓	✓	✓	✓	✓
Kibertér	✓	✓	✓		✓

2.6. ábra – MI-alkalmazások indikatív mátrixa katonai funkcionális területeken (saját szerkesztés; NATO és szakirodalom alapján [14], [18], [19]).

2.3.1. ISR és OSINT: össz-adatforrású feldolgozás és fúzió

Az ISR (Intelligence, Surveillance and Reconnaissance) terület az MI egyik legérettebb katonai alkalmazási doménje, mert a szenzoradatok mennyisége és komplexitása az emberi feldolgozó kapacitást régóta meghaladja. A DL-alapú kép- és videófeldolgozás (automatikus objektumfelismerés, követés, változásdetektálás) mellett különösen a össz-adatforrású (multi-INT) fúzió és az OSINT szerepe nő.

A Nemzetbiztonsági Szemlében megjelent OSINT-tanulmányom a nyílt információszerzés (OSINT) evolúcióját a „big data” kontextusában tárgyalja, és kiemeli a nagy adathalmazok jellemző dimenzióit (volume, velocity, variety, veracity, value – „5V”) mint az automatizálást kikényszerítő tényezőket [15, pp. 30–33]. A katonai tervezésben az OSINT nem önálló „INT” világként jelenik meg, hanem olyan, gyorsan elérhető, gyakran nagy területi lefedettségű inputként, amely a helyzetértékelést (situational awareness), a civil környezet (civil considerations) és az ellenség információs aktivitásának értelmezését kiegészíti.

A MI-vel támogatott OSINT feldolgozás tipikus lánc a források gyűjtésétől a tartalmi kivonatoláson át a fúzióig és az elemzői felülvizsgálatig terjed. A lánc kritikus pontjai katonai kontextusban a megbízhatóság (veracity) és a manipuláció: a nyílt források ellenfelek által célzottan torzíthatók, automatizált bot-hálózatokkal vagy generatív tartalmakkal (deepfake) mérgezhetők, ezért a fúzióban a trianguláció (több egymástól független forrás) és a forráskritika elengedhetetlen [15, pp. 40–46].

A NATO EDT-összefoglalója külön hangsúlyozza, hogy „az adat kulcs-enabler minden EDT számára”, és a Szövetség adat-exploitatív politikával és digitális transzformációs implementációs stratégiával erősíti a data-driven működést [14]. ISR/OSINT oldalról ez azt jelenti, hogy a tervezéstámogató MI-komponensek nem elszigetelt kísérleti megoldásokként, hanem adat-architektúrába ágyazva értelmezhetők: metaadat-szabványok, hozzáférés-kezelés, adatminőségi metrikák és auditálható feldolgozási láncok szükségesek a műveleti felhasználáshoz.

Az ISR/OSINT MI-alkalmazások egyik konkrét tervezési haszna a „fókuszált figyelem” biztosítása: a törzsek a COPD mission analysis és IPB/JIPOE folyamatában gyakran küzdenek azzal, hogy az információszerzés és az értékelés „szétterül” a túl sok lehetséges kérdés között. A gépi szűrés (prioritizálás), az automatikus témaklaszterezés, valamint a trend- és anomáliaajelzés képes az elemzői figyelmet a legnagyobb döntési relevanciájú jelenségekre irányítani – ugyanakkor a „black box” jellegű modellek esetén a döntési felelősség nem delegálható teljesen algoritmusokra (erről részletesen a 2.4–2.5 fejezetben).

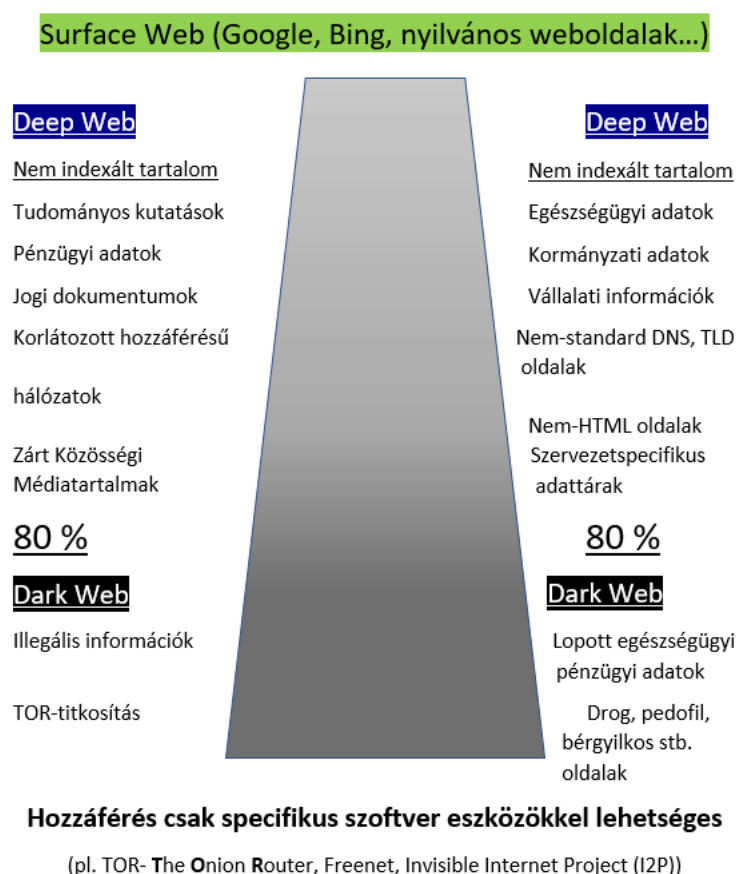
Aktív OSINT-tevékenység az avatárok használata, amelyek lényegében virtuális HUMINT-ügynököknek tekinthetők. Az avatár lehet például egy hamis személyazonossággal létrehozott profil, amely a célszemély valamely érdeklődési területét (politika, szex, gasztronómia, stb.) kihasználva lép kapcsolatba vele, kerül az ismeretségi körébe, férkőzik a bizalmába. Az avatárok lehetővé teszik az elemzők számára, hogy biztonságosan kommunikáljanak, figyeljék és manipulálják a célpontokat. Az avatárok a rendszerben gyorsan definiálhatók, módosíthatók vagy új attribútumok adhatók meg az egyes avatárokhoz.

Fentiekén túl fontos tisztázni a deep web – dark web környezetet. A deep web minden olyan internetes tartalomra vonatkozik, amelyet különböző okok miatt nem indexelnek olyan keresőmotorok, mint a Google. Ez a definíció tehát magában foglalja a dinamikus weboldalakat, a blokkolt webhelyeket (például azokat, amelyek a hozzáféréshez CAPTCHA választadást kérnek), a privát webhelyeket (például azokat, amelyekhez bejelentkezési hitelesítés szükséges), a nem HTML/kontextuális/szkriptes tartalmakat és korlátozott

hozzáférésű hálózatokat. A teljes deep web becslések szerint az internetes tartalmak mintegy 80%-át jelenti.

A korlátozott hozzáférésű hálózatok lefedik mindazokat az erőforrásokat és szolgáltatásokat, amelyek normál esetben nem elérhetők a szabványos hálózati konfigurációval, így lehetőséget kínálnak a rosszindulatú szereplők fellépésére. Idetartoznak azok a webhelyek, amelyek domain-nevei olyan Domain Name System (DNS) rootokon vannak regisztrálva, amelyeket nem az Internet Corporation for Assigned Names and Numbers (ICANN) kezel, és ezért olyan nem-szabványos felső szintű domaineikkel (*Top-Level Domains*, TLD) rendelkező URL-eket tartalmaznak, amelyekhez egy adott DNS-kiszolgálóra van szükség.

További példa a domain-nevüket a szabványos DNS-től teljesen eltérő rendszeren regisztráló webhelyek, mint például a .BIT tartományon regisztrált „Bitcoin Domain”. Ezek a rendszerek az ICANN által előírt domainnév-szabályozást kikerülik, az alternatív DNS-ek decentralizált jellege szintén nagyon megnehezíti ezeknek a tartományoknak a beazonosítását.



2.7. ábra: Az Internet felosztása: surface web – deep web – dark web. Forrás: a szerző szerkesztése

A korlátozott hozzáférésű hálózatok között találhatók a darknet-hálózatok, olyan infrastruktúrákon tárolt webhelyek, amelyek speciális szoftverek – például a TOR – használatát teszik szükségessé az elérésükhöz. A dark web és a deep web között különbség van, bár egyesek szinonimaként használják őket, de a dark web csak egy része a deep webnek. A dark web hálózatokon titkosított peer-to-peer kapcsolat jön létre a partnerek között. A dark web rendszerek példái közé tartozik a TOR, a Freenet vagy az Invisible Internet Project (I2P). A dark web a törvénytelen cselekményekkel kapcsolatos információk adattárháza. Idetartoznak a malware-szoftverekkel kapcsolatos információktól a pedofil oldalakon át a bérgyilkosságot végrehajtó hirdetésekig szinte minden, ami törvénytelen. Természetesen mind a deep webből, mind a dark webből kinyerhetők OSINT-módszerekkel, a megfelelő eszközökkel információk.

2.3.2. C2 és tervezéstámogatás: COA-generálás, wargaming, optimalizáció

A C2 (Command and Control) és a műveleti tervezés támogatása a disszertáció fókuszterülete. Itt az MI-alkalmazások nemcsak a harcmezőn megjelenő autonóm rendszerekben, hanem a törzsszintű döntés-előkészítésben jelennek meg: a helyzetértékelés strukturálásában, a COA-k előállításában, a kockázatok és erőforrásigények becslésében, valamint a végrehajtási változatok összehasonlításában.

A Négyesi–Szabadszabó tanulmány az MI-támogatott katonai döntéstámogató rendszerekben több módszert említ, különös tekintettel azokra, amelyek a bizonytalanság kezelésére és az optimalizációra alkalmasak: Bayes-féle megközelítések, fuzzy logika, genetikai algoritmusok, szakértői rendszerek és neurális hálók [16, p. 30]. E módszerek jól illeszthetők a COPD azon fázisaihoz, ahol (a) több alternatíva összevetése szükséges (COA analysis), (b) erőforráskorlátos problémák merülnek fel (force allocation, logistics), vagy (c) a bizonytalanság domináns (intel gaps, ellenség szándék).

A tervezéstámogató MI egyik tipikus felhasználása a COA-generálás és -variálás. Itt a rendszer nem „dönt” a parancsnok helyett, hanem alternatívák térképét (option space) bővíti: sablonokból, múltbeli tapasztalatokból, szimulációs eredményekből és korlátokból kiindulva lehetséges műveleti változatokat javasol. A klasszikus, szabályalapú COA-generálás előnye az átláthatóság; a generatív MI (pl. LLM) előnye, hogy képes gyorsan strukturált szöveges vázlatokat készíteni (pl. döntési pontok, feladatok, feltételezések). A katonai alkalmazás azonban csak akkor tekinthető érettnek, ha a generált változatok ellenőrizhetők, a források és adatok auditálhatók, és a biztonsági kockázatok kezeltek [16], [10,p.12–17].

A wargaming és szimulációs támogatás terén az MI két szinten jelenik meg: (1) ellenség viselkedésének modellezése (pl. RL-alapú agentek), (2) gyors kiértékelés és érzékenységi elemzés. A katonai tervezésben a wargaming célja a COA-k robusztusságának tesztelése; az MI itt a futtatások számának növelésével és a kimenetek automatikus összegzésével csökkentheti az emberi elemzők terhelését.

Végül kiemelendő az optimalizáció: a stratégiai–műveleti tervezés gyakran többcélú optimalizáció, ahol a célok (pl. idő, kockázat, erőforrás, politikai költség) ütköznek. A genetikus algoritmusok és más kereső-heurisztikák segíthetnek nagy kombinatorikus terekben (pl. ütemezés, logisztikai hálózatok), míg a fuzzy logika és Bayes-féle modellek a bizonytalanság formalizálásában támogatják a döntéshozót [16, p. 30].

2.3.3. Elektronikai hadviselés (EW): spektrum, jelazonosítás, adaptív zavarás

Az elektronikai hadviselés (EW) a spektrumbeli környezet gyors változásai és az emissziók heterogenitása miatt szintén erősen adat- és mintázatintenzív terület. A NATO STO jelentés több helyen utal arra, hogy az EDT-k azonosításakor a C4ISR-re és a túlélőképességre gyakorolt hatás központi szempont [2, p. 9]. Az EW-ben a gépi tanulás a következő feladatokban jelenik meg:

- Jelazonosítás és osztályozás (RF fingerprinting): ismeretlen emissziók klaszterezése, platform-azonosítás;
- Spektrumhelyzet-kép kialakítása: valós idejű spektrumhasználat térképezése, interferenciák detektálása;
- Adaptív zavarás és védelem: olyan „cognitive EW” logikák, ahol a rendszer a környezetből tanulva választ zavarási vagy elhárítási paramétereket.

A tervezéstámogatásban az EW MI-alapú eszközei a fenyegetési modell és a kommunikációs biztosítás (comms assurance) rétegét erősítik: a COA-k értékelése során a spektrumbeli sebezhetőségek, a zavarási kockázatok és az alternatív kommunikációs útvonalak (pl. next-gen comm networks) összevetése egyre inkább számításigényes feladat. A NATO EDT-listája külön kezeli a következő generációs kommunikációs hálózatokat mint kiemelt területet [14], ami jelzi, hogy a spektrum és kommunikációs infrastruktúra a tervezésben is stratégiai tényezővé vált.

2.3.4. Logisztika és fenntartás: predikció, ellátási hálózatok és erőforrás-optimalizáció

A logisztika és fenntartás a katonai műveletek „gerince”, és a NATO/partneri környezetben különösen összetett, mert többnemzeti ellátási láncokat, eltérő szabványokat és sokszor korlátozott infrastruktúrát kell kezelni. A NATO ACT MUAAR kezdeményezés is kiemeli a logisztikát mint olyan területet, ahol az AA&R jelentős hatékonysági és költségmegtakarítási potenciált hordoz [7, p. 1].

A Land Forces Academy Review-ban publikált cikkem összefoglalója szerint az MI a katonai logisztikában többek között prediktív karbantartásban, készletgazdálkodásban, útvonal- és szállítás-optimalizációban, valamint a veszteségek előrejelzésében alkalmazható [14, pp. 161–164]. A tervezésben ezek a képességek akkor hasznosulnak, ha a COA értékelés részeként valószínűsíthető ellátási kockázatokat, fogyási ütemeket és „bottleneck”-eket lehet becsülni. A logisztikai MI-alkalmazások sajátossága, hogy gyakran strukturált adatbázisokon (készlet, igény, útvonal, kapacitás) futnak, így a magyarázhatóság és az auditálhatóság technikailag könnyebben biztosítható, mint a tisztán DL-alapú, érzékelési feladatokban. Ugyanakkor a valós idejű adatok (IoT szenzorok, járműtelemetria) és a többnemzeti adatmegosztás biztonsági kockázatai miatt a kiberbiztonsági és adatkezelési követelmények itt is kritikusak – ezt a 2.4 alfejezet tárgyalja részletesen.

2.3.5. Kiber (Cyber): anomáliadetektálás, fenyegetés-intelligencia, aktív védelem

A kiber műveletekben az MI szerepe kettős: (1) a támadók is alkalmazhatnak MI-t a sebezhetőségek felderítésére, a social engineering támogatására vagy automatizált exploit-láncok építésére; (2) a védelemben a nagy volumenű hálózati és rendszerlogok feldolgozása, az anomáliadetektálás és a fenyegetés-felderítés (Threat Intelligence) integrációja MI nélkül egyre kevésbé skálázható.

Tervezési oldalon a kiber dimenzió abban válik hangsúlyossá, hogy a műveleti környezet részeként (OE) a kritikus infrastruktúra, a kommunikációs rendszerek és a digitális szolgáltatások sebezhetőségei közvetlenül befolyásolják a COA-k végrehajthatóságát és kockázati profilját. A NATO EDT-stratégiája a „protect” logikában kifejezetten kiemeli, hogy a Szövetség védi innovációs ökoszisztémáját az ellenséges beavatkozástól [14]; ez a kiberbiztonságot nemcsak technikai, hanem stratégiai kérdéssé is emeli.

A kiber MI-megoldások tipikusan (a) felügyelt osztályozási (malware family), (b) felügyelet nélküli anomáliadetektálási (outlier) és (c) graf-alapú (kapcsolati) módszereket használnak. A

tervezéstámogatásban ezek az eszközök képesek a fenyegetési képet gyorsabban frissíteni és a kiberhatások becslését megalapozni (pl. milyen valószínűséggel zavarható meg egy C2-hálózat, milyen redundancia szükséges).

2.3.6. Összegző megállapítások

Az ISR/OSINT, C2/tervezéstámogatás, EW, logisztika és kiber területeken bemutatott példák közös tanulsága, hogy az MI értéke a katonai környezetben elsősorban a (1) nagy adathalmazok értelmezésének gyorsításában, (2) alternatívák gyors előállításában és összevetésében, valamint (3) a komplex rendszerek optimalizációjában jelenik meg. A NATO EDT-keretei és az AI/autonómia stratégiák ugyanakkor hangsúlyozzák: az adaptáció gyorsítása csak akkor lehet fenntartható, ha az MI alkalmazások biztonságosak, megbízhatóak és a felelős használat elveit követik [14], [16], [20].

Ennek megfelelően a következő alfejezetek (2.4–2.5) a katonai alkalmazhatóság kritériumait tárgyalják: a megbízhatóság, magyarázhatóság, adatminőség és kiberbiztonság követelményeit, illetve az emberi kontroll (human-in-the-loop / human-on-the-loop) és az autonómia viszonyát. Ezek a szempontok képezik majd a COPD-fázisokhoz illesztett MI-támogatási modell értékelési keretrendszerét is.

2.4. Megbízhatóság, magyarázhatóság, adatminőség és kiberbiztonsági követelmények

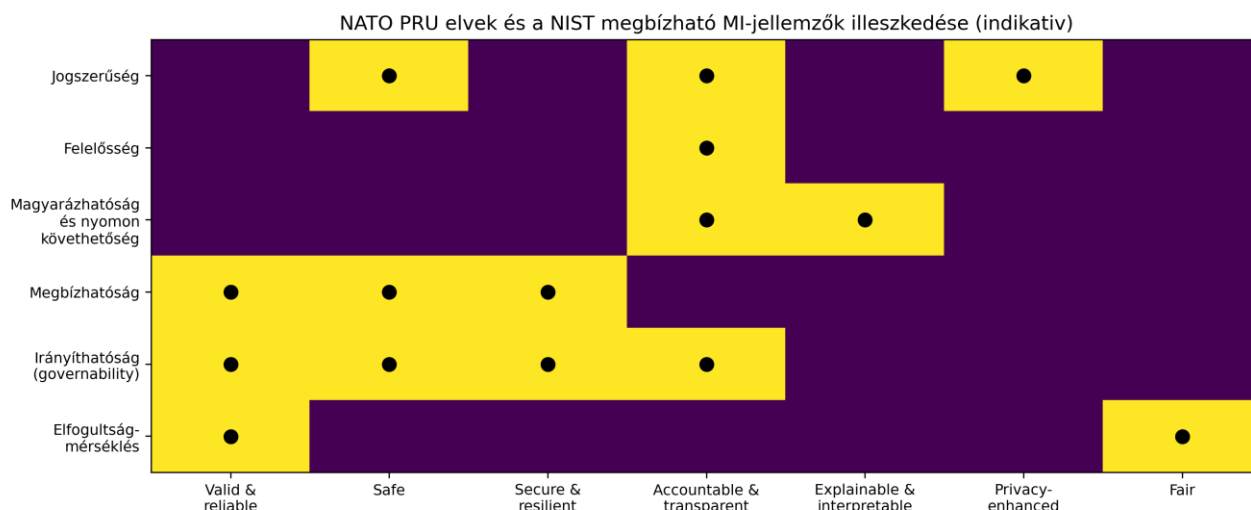
A katonai környezetben alkalmazott MI-rendszerek – különösen a döntéstámogatást, hírszerzési feldolgozást (ISR/OSINT), C2-t és a műveleti tervezést támogató megoldások – nem pusztán „teljesítmény” (pontosság, gyorsaság) szerint értékelendők. A műveleti kockázat, a koalíciós interoperabilitás, az adat- és rendszerbiztonsági kitettség, valamint az elszámoltathatóság együttesen teszik szükségessé az ún. megbízható/trösztölhető (trustworthy) MI-követelményrendszer alkalmazását. A NATO a felelős alkalmazást a Principles of Responsible Use (PRU) elvein keresztül operacionalizálja [16], [15], míg a polgári–ipari szférában a NIST AI Risk Management Framework (AI RMF 1.0) széles körben hivatkozott keretet ad a kockázatalapú megközelítéshez [12]. A két megközelítés közös metszete releváns a Magyar Honvédség (MH) és a NATO/EU műveleti környezetében is, mert a koalíciós

műveletekben az MI-rendszerek bizalmi szintje és „hitelesíthetősége” közvetlenül befolyásolja a döntéshozatali ciklus minőségét és a műveleti kimenetet.

2.4.1. Megbízhatóság és robusztusság: „pontosság” helyett életciklus-szemlélet

A megbízhatóság (reliability) katonai értelmezésben nem egyetlen metrika, hanem egy életcikluson át biztosítandó tulajdonsághalmaz: (i) validitás és reprodukálhatóság (a modell a definiált feladatra alkalmas és az eredmények megismételhetők); (ii) robusztusság (zajos, hiányos, ellenségesen manipulált vagy eltolódó adatkörnyezetben is elfogadható teljesítmény); (iii) biztonság és rendszerbiztonság (safety) – különösen emberéletet, kritikus infrastruktúrát vagy erőalkalmazást érintő döntéseknél; (iv) kiberreziliencia (a modell és a teljes MLOps-lánc védelme); valamint (v) működtethetőség, fenntarthatóság, monitorozhatóság. A NIST AI RMF a „valid and reliable”, „safe”, „secure and resilient” jellemzőket együtt kezeli a megbízható MI magját adó karakterisztikákként [12]. Ennek katonai megfelelője a NATO PRU „reliability” elve, amelyhez kapcsolódóan a NATO irányelvek a minőségbiztosítás és a kockázatkezelés beépítését hangsúlyozzák a fejlesztés–tesztelés–bevezetés teljes ciklusába [16], [15].

A „modellek érettsége” katonai környezetben nem kizárólag TRL-szint, hanem VVTE-érettség (verification, validation, test & evaluation), adat-érettség (adatgazdálkodás, hozzáférés, címkézés, metaadatok), valamint üzemeltetési érettség (monitoring, incidenskezelés, frissítések, konfigurációkezelés). A 2021-es amerikai védelmi minisztériumi (DoD) RAI memorandum a „Reliable” elvet explicit módon a teljes életciklusra kiterjedő tesztelési és assurance követelményként fogalmazza meg [21]. A gyakorlati következmény: a katonai MI-megoldásoknál már a követelményrendszerben (requirements engineering) rögzíteni kell a használati határokat, a tiltott felhasználási módokat, a környezeti feltételezéseket (assumptions) és a „degradációs viselkedést” (mit csinál a rendszer, ha az adat romlik, a szenzor kiesik, vagy a kommunikáció korlátozott).



2-5. ábra – NATO PRU elvek és a NIST AI RMF megbízható MI-jellemzők indikativ megfeleltetése (saját szerkesztés [16], [12] alapján).

2.4.2. Magyarázhatóság, nyomon követhetőség és auditálhatóság

A magyarázhatóság (explainability) és a nyomon követhetőség (traceability) katonai környezetben két okból kritikus: (1) a döntéshozói felelősség és elszámoltathatóság (accountability) – különösen, ha az MI kimenete közvetve vagy közvetlenül befolyásolja az erőalkalmazást, célkijelölést, vagy a műveleti kockázatvállalást; (2) a koalíciós bizalom és interoperabilitás – a partnernemzetek és parancsnoki szintek számára értelmezhetővé kell tenni, hogy „miért ezt javasolja” a rendszer. A Szabadszöldi által elemzett katonai MI-alkalmazások egyik visszatérő korlátja éppen a mélytanuló rendszerek „black box” jellege, amely a szakértői validációt és a felelősségi láncot nehezíti [10, pp. 158, 164–165].

A DoD AI Ethical Principles közül a „Traceable” elv az átlátható és auditálható módszertanok, adatforrások és dokumentáció követelményét rögzíti [21]. Hasonlóan, az EU nagy szakértői csoportja (HLEG) a „Transparency” és az „Explicability” elvet a megbízható MI egyik alapköveként emeli ki, és a dokumentáltságot, kommunikálhatóságot, valamint a döntési folyamatok visszakövethetőségét hangsúlyozza [22].

Katonai döntéstámogatásban a magyarázhatóság gyakorlati megoldása több szinten valósítható meg: (i) modell-szintű magyarázat (pl. feature importance, lokális magyarázatok, saliency map); (ii) adat-szintű transzparencia (adatforrások, frissesség, bizonytalanság és minőségi címkék megjelenítése); (iii) folyamat-szintű nyomon követés (MLOps naplózás, verziókezelés, döntési audit trail); (iv) felhasználói felület (C2/tervezői környezet) szintjén magyarázó

narratívák, ellenőrző kérdések és alternatívák. A cél nem az, hogy a döntéshozó „levezesse” a neurális hálózat belső működését, hanem hogy megkapja a döntéshez szükséges bizonyosságot: milyen adatból, milyen feltételek mellett, milyen bizonytalansággal született a javaslat, és hogyan ellenőrizhető.

2.4.3. Adatminőség és adatgazdálkodás: a katonai MI szűk keresztmetszete

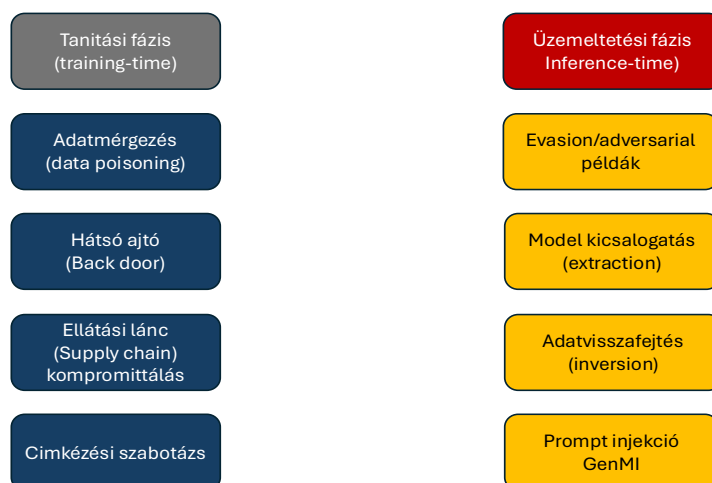
Az MI-rendszerek teljesítményét és megbízhatóságát alapvetően determinálja az adatminőség. A katonai környezetben az adat jellemzően heterogén (több szenzoros ISR, HUMINT, SIGINT, OSINT, logisztikai és személyzeti adatok), részben minősített, időben gyorsan avul, és ellenséges információs műveleteknek kitett. Szabadföldi OSINT-elemzése a Big Data „5V” keretrendszerét emeli ki (volume, velocity, variety, veracity, value), ahol a veracity – az információ hitelessége és manipulálhatósága – a katonai felhasználásban különösen kockázatos dimenzió [11, p. 30]. Ugyanez a szerző a neurális hálózatok belső reprezentációinak nehezebb értelmezhetőségét és a hibák „láthatatlanságát” is hangsúlyozza, ami adatminőségi problémákkal kombinálva a döntéstámogatásban torzításokat eredményezhet [11, p. 35].

A NATO AI-stratégiai dokumentumai következetesen a „good data” előfeltételére építenek: a felelős használat és a skálázható bevezetés feltétele a kormányzott, megbízható és védett adat-ökoszisztéma [23], [15]. A gyakorlatban ez – különösen koalíciós környezetben – adatstandardok, metaadat-sémák, minőségi címkék, megosztási policyk és adat-hozzáférési mechanizmusok egységesítését jelenti. Katonai döntéstámogatásban kiemelt követelmény az adat „provenance” (eredet) és „lineage” (feldolgozási lánc) kezelése: a tervezőnek tudnia kell, milyen forrásból származik az információ, milyen feldolgozási lépések történtek, és hol lehetett a lánc sérülékeny (pl. OSINT forrás manipuláció, deepfake, automatizált bot-hálózat).

Az adatminőség biztosításának tipikus katonai kontrolljai: (i) adatforrások minősítése és „bizalmi szint” hozzárendelése (trusted/untrusted), (ii) adatvalidációs szabályok és anomália-detektálás a beérkezéskor, (iii) címkézési protokollok (human label + second reviewer, mérési hibák jelölése), (iv) torzítás- és reprezentativitás-analízis (bias), (v) adatfrissesség és „drift” monitorozás, (vi) minősített adatok esetén hozzáférés- és jogosultságkezelés (need-to-know), valamint (vii) koalíciós adatmegosztásnál a releasability és sanitization szabályok. A cél: az MI ne „minden áron” adjon választ, hanem a bizonytalanságot és a forráskockázatot is láthatóvá tegye.

2.4.4. Kiberbiztonság és adversarial fenyegetések: MI mint támadási felület

Az MI-rendszerek kiberbiztonsága nem azonos a hagyományos IT-rendszerek védelmével: a fenyegetések jelentős része modell- és adat-specifikus (adversarial ML). A NIST adversarial ML taxonómiája külön kezeli a tanítási fázisban végrehajtott támadásokat (pl. data poisoning, backdoor) és az inference fázisban végrehajtott támadásokat (pl. evasion, model extraction, model inversion) [24]. A katonai környezetben mindehhez hozzáadódik az ellátási lánc kockázata (harmadik féltől származó adat, előtanított modell, szoftverkomponens), valamint az ellenséges fél célzott megtévesztése (deception), például adversarial példák vagy deepfake tartalmak alkalmazásával.



Megjegyzés: A fenyegetések taxonómiája és a mitigáció a NIST AI 100-2e és a MITRE ATLAS alapján rendszerezhetők

2-7. ábra – MI/ML rendszerek tipikus támadási felülete (saját szerkesztés; NIST adversarial ML taxonómia és MITRE ATLAS alapján [24], [25]).

A DoD Directive 3000.09 (Autonomy in Weapon Systems) a fegyverrendszer-autonómia vonatkozásában a „failure” okai közé explicit módon beemeli az ellenséges kiber- és ellátási lánc támadásokat, a spoofing/jamming jelenségeket és az előre nem látott harctéri helyzeteket is, és ezek minimalizálását a tervezés–tesztelés–üzemeltetés egészére kiterjedő követelményként kezeli [19]. Bár a műveleti tervezést támogató MI-rendszerek nem fegyverrendszerek, a kiberreziliencia logikája azonos: a tervezői környezetben használt modellek (pl. helyzetkép-fúzió, logisztikai előrejelzés, COA-értékelés) elleni támadás a parancsnoki döntéshozatalt támadja, így a kiberkockázat közvetlen műveleti kockázattá válik.

A védelem gyakorlati eszköztára kiterjed: (i) MLSecOps/DevSecOps integráció (a biztonság „beépítése” a modell életciklusába), (ii) adat- és modell-integritás ellenőrzések (hash, SBOM/MBOM jellegű komponens-leltár, szignált modellek), (iii) hozzáféréskezelés és titkosítás, (iv) adversarial tesztelés (red teaming), (v) kimenet-monitoring és „guardrail” mechanizmusok különösen generatív MI-nél (prompt-injekció, adat-kiszivárgás), valamint (vi) incidenskezelési protokollok. Az OpenSSF MLSecOps whitepaper kiemeli, hogy a klasszikus DevSecOps mintázatok csak részben elégségesek, mert az ML rendszerek dinamikusak, adatintenzívek és új típusú fenyegetéseket hordoznak; ezért a biztonsági baseline, kockázatfelmérés és folyamatos monitoring a teljes MLOps-láncban szükséges [26].

2.4.5. VVTE, minőségbiztosítás és folyamatos monitorozás (MLOps)

A katonai MI-rendszerek validációja és hitelesítése (VVTE) akkor tekinthető megfelelőnek, ha a tesztelés nem csak „labor” környezetben, hanem releváns műveleti szcenáriókban (operationally relevant scenarios) és stresszhelyzetekben is megtörténik. Ennek része: (i) funkcionális teszt (helyes működés), (ii) robusztussági teszt (zaj, hiányosság, adatelcsúszás), (iii) adversarial teszt (célzott megtévesztés), (iv) kiberbiztonsági teszt (komponensek, API-k, jogosultságok), (v) felhasználói teszt (ember–gép interakció, UI/UX a C2/tervezői környezetben), és (vi) koalíciós teszt (interoperabilitás, adatmegosztás, policy kompatibilitás).

A DoD RAI governance-szemlélete (RAI Governance, Warfighter Trust, System Engineering, Ecosystem) arra utal, hogy az MI minősége nem „projektvégi” ellenőrzés, hanem folyamatos szervezeti és műszaki kontroll [21]. A NIST AI RMF pedig a GOVERN, MAP, MEASURE, MANAGE funkciók mentén állít fel kockázatkezelési életciklust, amelyben a mérések (performance + risk) és a menedzsment intézkedések folyamatosan visszacsatolnak [12]. Katonai tervezéstámogatásban ez különösen fontos, mert a műveleti környezet változásai (ellenség taktikaváltás, infrastruktúra romlása, időjárás, civil mintázatok) gyorsan „driftet” okozhatnak, amit a rendszernek észlelnie kell.

MLOps/MLSecOps életciklus és tipikus kontrollpontok (saját szerkesztés)



Kontrollpontok (példa): adatminőség, hozzáférés- és jogosultságkezelés, XAI/VVTE, kiberbiztonság, drift-detektálás, incident response

2-6. ábra – MLOps/MLSecOps életciklus és kontrollpontok (saját szerkesztés; NIST AI RMF és OpenSSF MLSecOps megfontolások alapján [12], [26]).

2.4.6. Követelményrendszer-összegzés: minimum-követelmények katonai MI-hez

Az előző alfejezetek alapján – a NATO PRU és a NIST/DoD keretek közös metszetére építve – a katonai MI-rendszerekkel szemben az alábbi minimum-követelmények fogalmazhatók meg (a disszertáció későbbi fejezeteiben ezek COPD-fázisokra és tervezési termékekre lesznek leképezve):

- Definiált rendeltetés és határok: a rendszer feladata, tiltott használatai, feltételezései és degradációs viselkedése dokumentált.
- Adatgazdálkodás: adatprovenance, minőségi címkék, hozzáférés- és megosztási policyk, torzításvizsgálat és frissesség-kezelés.
- Magyarázhatóság és audit: döntési naplózás, verziókezelés, audit trail; a felhasználói felületen bizonytalanság és forráskockázat megjelenítése.
- VVTE és red teaming: robusztusság- és adversarial tesztek releváns scénáriókban; eredmények dokumentálása, ismételhetőség.
- Kiberbiztonság/MLSecOps: ellátási lánc kontroll, hozzáférés, integritás, incidenskezelés; GenMI-nél guardrail mechanizmusok.
- Emberi kontroll: HITL/HOTL működés, felelősségi rend és beavatkozási jogok; képzés és eljárásrend.

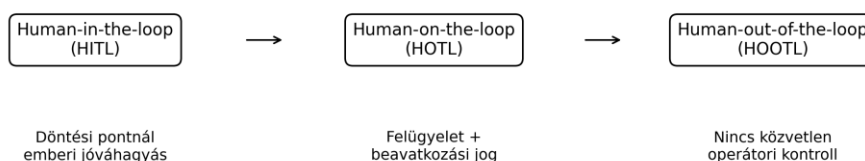
2.5. MI és autonómia kapcsolata; emberi kontroll (human-in-the-loop/on-the-loop)

Az MI és az autonómia viszonya a katonai alkalmazások egyik legvitatottabb, ugyanakkor leggyorsabban fejlődő területe. Fontos megkülönböztetni (i) az automatizálást (előre definiált szabályok mentén), (ii) az MI-alapú adaptív automatizálást (adatvezérelt döntési szabályok), és (iii) az autonómiát, amikor a rendszer a környezet érzékelése és értelmezése alapján önállóan választ cselekvési opciót. A NATO az autonómia bevezetését külön „Autonomy Implementation Plan” mentén támogatja, amelynek célja, hogy a képességek skálázott bevezetése mellett a felelős használat is érvényesüljön [20].

2.5.1. Fogalmi keret: autonóm, fél-autonóm és operátor-felügyelt rendszerek

A DoD Directive 3000.09 definíciós keretet ad az autonóm, fél-autonóm és operátor-felügyelt autonóm fegyverrendszerekre. A direktíva szerint az autonóm fegyverrendszer „aktiválás után” képes célokat kiválasztani és leküzdeni további operátori beavatkozás nélkül; az operátor-felügyelt autonóm rendszer viszont beavatkozási és megszakítási lehetőséget biztosít a kezelőnek még elfogadhatatlan károkozás bekövetkezése előtt [19]. A definíciós keret a döntéstámogató rendszerekre is adaptálható: az autonómia „tárgya” itt nem feltétlen a fegyverhasználat, hanem például a tervezési termékek generálása, COA-k rangsorolása, logisztikai útvonal-optimalizálás vagy hírszerzési jelzések prioritizálása.

Emberi kontroll-modellek és autonómia: alapfogalmi áttekintés (saját szerkesztés)



A műveleti tervezést támogató MI-nél tipikusan HITL/HOTL megközelítés indokolt; fegyverrendszer-autonómiánál a kontroll- és VVTE-követelmények a legszigorúbbak.

2-8. ábra – Emberi kontroll-modellek (HITL/HOTL/HOOTL) és alkalmazhatóságuk (saját szerkesztés [19] alapján).

2.5.2. Emberi kontroll: felelősség, beavatkozási jog és „meaningful human control”

Az emberi kontroll katonai MI-alkalmazásban nem pusztán technikai, hanem doktrinális és jogi kérdés. A kontroll megtervezésekor három dimenziót kell egyszerre kezelni: (i) idő (van-e elegendő idő emberi jóváhagyásra), (ii) információ (az ember érti-e a rendszer javaslatának alapját és bizonytalanságát), valamint (iii) beavatkozási képesség (képes-e a rendszer működését megállítani, felülbírálni vagy korlátozni). A DoD AI Ethical Principles „Responsible” és „Governable” elvei explicit módon rögzítik, hogy a felelősség a DoD személyzeté marad, és a rendszereket úgy kell tervezni, hogy nem kívánt viselkedés esetén lekapcsolhatók legyenek [21]. A NATO PRU keret ugyanezt a felelősségi és irányíthatósági logikát követi [16], [15].

A műveleti tervezést támogató MI tipikus „helye” a HITL/HOTL spektrumon a döntéstámogatási ciklus különböző pontjain változhat. Például: adat-előkészítés és korai helyzetértékelés (JIPOE jellegű feladatok) esetén HOTL modellt működhet, ahol a rendszer automatikusan priorizál, de az elemző felülvizsgál; COA-generálásnál HITL célszerű, ahol a rendszer alternatívákat javasol, de a parancsnoki döntéshozó választ; kritikus kockázatú „automatikus” ajánlásoknál pedig a rendszernek képesnek kell lennie bizonytalanság esetén „visszaadni” a döntést (graceful degradation).

2.5.3. Autonómia, etikusság és jogi megfelelés: következmények a tervezéstámogatásra

Az autonóm rendszerek körüli etikai és jogi diskurzus a fegyverrendszereken keresztül vált hangsúlyossá, de a döntéstámogató MI-re is kihat. Egyrészt, a döntéshozói felelősség nem delegálható „algoritmusra”; másrészt, a hibás vagy manipulált MI-kimenet a parancsnoki döntésen keresztül műveleti és – végső soron – jogi következményekhez vezethet. A Négyesi–Szabadföldi tanulmány kifejezetten rámutat arra, hogy a DoD öt etikai elve (Responsible, Equitable, Traceable, Reliable, Governable) a katonai MI-bevezetés gyakorlati keretrendszerét adja, és a tervezéstámogatásban is alkalmazandó [12, p. 33]. A NATO és szövetséges rendszerekben ez a keret a PRU elvekben jelenik meg [16], [15].

A disszertáció szempontjából lényeges következtetés: a COPD szerinti műveleti tervezésben az MI-t úgy kell integrálni, hogy (i) a felelősségi lánc egyértelmű maradjon (ki dönt, ki ellenőriz), (ii) a tervezési termékek auditálhatók legyenek (adatforrás, változások, verziók), és (iii) a koalíciós környezetben a „bizalom” transzparens kontrollokkal legyen megalapozva.

E követelmények közvetlenül befolyásolják az MI-támogatás tervezési „feladatait” és a későbbi fejezetekben bemutatott alkalmazási mintázatokat.

2.6. Összefoglalás

A 2. fejezet (2.1–2.5) elméleti kerete alapján a műveleti tervezést támogató MI katonai alkalmazása akkor tekinthető megvalósíthatónak és skálázhatónak, ha a technológiai teljesítményen túl a megbízhatóság, magyarázhatóság, adatminőség és kiberbiztonság követelményei is teljesülnek. A NATO PRU elvek, a NIST kockázatkezelési megközelítése és a DoD RAI keretei konzisztens módon azt sugallják, hogy a „trust” a governance–VVTE–MLOps–emberi kontroll négyesére épül. A fejezet külön hangsúlyozta: (i) a döntéstámogató MI támadási felületet képez (adversarial ML), (ii) az adatminőség a katonai MI legfontosabb szűk keresztmetszete, és (iii) az autonómia kérdése az emberi kontroll és felelősség újragondolását igényli. A következő fejezetekben a disszertáció ezeket a követelményeket a COPD módszertan fázisaira és tervezési termékeire vetíti, és konkrét alkalmazási lehetőségeket, valamint értékelési keretrendszert dolgoz ki.

3. A KATONAI MŰVELETTERVEZÉS FEJLŐDÉSE ÉS ELMÉLETI ALAPJAI

A fejezet célja, hogy rövid, de rendszerezett áttekintést adjon a katonai művelettervezés fogalmi alapjairól és történeti fejlődéséről. Az áttekintés a stratégiai–hadműveleti–harcászati szintek logikájára épül, majd kronologikus rendben bemutatja azokat a meghatározó fordulópontokat, amelyek a modern, többnemzeti és összhaderőnemi (joint) tervezési eljárásokhoz vezettek. A fejezet a későbbi, MI-támogatásra fókuszáló fejezetekhez biztosít elméleti keretet.

3.1. A művelettervezés fogalmi rendszere (stratégiai–hadműveleti–harcászati szint)

A katonai művelettervezés a célok (ends), a módszerek (ways) és a rendelkezésre álló eszközök/erőforrások (means) összehangolásának folyamata; egyszerre épít formális eljárásokra és a parancsnok alkotó, intuitív döntéseire. A NATO-doktrína a művelettervezést úgy kezeli, mint „operations planning”-et, amely stratégiai, hadműveleti és harcászati szinten egyaránt megjelenik, és külön hangsúlyozza, hogy az „operational planning” kifejezés kerülendő, mert összekeverhető az operatív szintű tervezéssel [27]. Az amerikai összhaderőnemi doktrína szerint az operatív művészet (operational art) a parancsnokok és törzsek kognitív megközelítése a stratégiák, kampányok és műveletek kialakítására, az ends–ways–means és kockázat integrálásával [9].

A három klasszikus háborús szint funkcionális logikája a következőképpen ragadható meg (3.1. ábra): (1) stratégiai szinten történik a politikai célok és a katonai célrendszer összekapcsolása, valamint a fő erőforrások, prioritások és korlátozások meghatározása; (2) hadműveleti/operatív szinten zajlik a kampány- és hadjáratlogika kialakítása, a műveleti megközelítés (operational approach) és a koncepció (CONOPS) kidolgozása, majd a részletes terv (OPLAN) létrehozása; (3) harcászati szinten a kijelölt feladatok végrehajtása, a harcászati manőver, a tüzek, a támogatások és a csapatok közvetlen alkalmazása dominál.



Forrás: saját szerkesztés NATO AJP-5 és US JP 5-0 alapján.

3.1. ábra. A háború szintjei és a művelettervezés jellemző termékei

A modern doktrínák a tervezést nem pusztán „feladatlistaként”, hanem a hadműveleti probléma keretezését és a megoldási lehetőségek kialakítását támogató gondolkodási keretként kezelik. A NATO AJP-5 megfogalmazásában az operatív művészet az a kritikus kapocs, amely összeköti a stratégiát és a harcászatot, és meghatározza, hogy a rendelkezésre álló erők milyen cselekvéseket hajtsanak végre időben és térben a kívánt hatások és célok elérése érdekében [27]. A tervezés ezért tipikusan két, egymást kiegészítő részfolyamatban jelenik meg: (a) művelettervezés/design (konceptcionális rész) és (b) részletes tervezés/tervtermékek előállítása (procedurális rész).

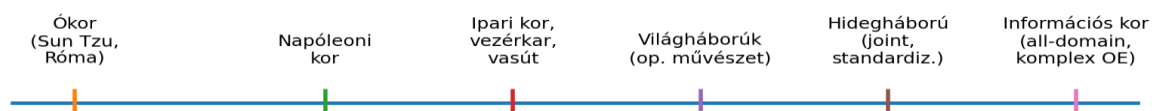
Szint	Fókusz	Jellemző kimenetek
Stratégiai	Célrendszer, erőforrások, korlátozások	Stratégiai iránymutatás, kampányirány, erőforrás-allokáció
Hadműveleti/operatív	Műveleti megközelítés és összhaderőnemi szinkron	CONOPS, OPLAN, LOE/LOO, ütemezés, fő erő- és képességigény
Harcászati	Kijelölt feladatok végrehajtása	Parancsok, harcászati tervek, végrehajtási eljárások

3.1. táblázat. A tervezés fókusza és jellemző kimenetei szintenként (saját szerkesztés).

3.2. Történeti áttekintés: ókortól a napóleoni háborúkig

A művelettervezés gyökerei az ókori hadtudományban és államszervezésben kereshetők: a hadjáratok sikere már ekkor is a célok, az erők, az időzítés, a felderítés és az utánpótlás összehangolásán múlt. Sun Tzu klasszikus tétele szerint „a hadviselés alapja a megtévesztés” [28], és a döntéshozatal egyik kulcsa az önismeret és az ellenség megismerése („ismerd az ellenséget és ismerd önmagad”) [28]. Ezek a gondolatok a mai tervezési ciklusokban (helyzetértékelés, ellenség-elemzés, kockázatkezelés) tovább élnek, noha a módszerek és eszközök nagyságrendekkel változtak.

Az ókori római hadszervezet és logisztikai rendszer (úthálózat, táborozási rend, standard eljárások) a „procedurális” tervezés korai mintáját adta: a nagyszámú csapat mozgatása és ellátása csak szabványosított folyamatokkal volt kezelhető. A középkorban és a kora újkorban a haderők professzionalizálódása, a tüzérség és a műszaki csapatok fejlődése erősítette a tervezés szerepét, de a döntő ugrást a napóleoni korszak hozta: a tömeghadseregek, a hadtest-rendszer és a gyors manőver igénye olyan „középszintű” (stratégia és harcászat közötti) koordinációt kényszerített ki, amelyet a szakirodalom a modern operatív gondolkodás előfutáraként kezel [29].



Forrás: saját szerkesztés.

3.2. ábra. A művelettervezés fejlődésének főbb korszakai (vázlatos idővonal)

A napóleoni háborúk tapasztalatai egyúttal ráirányították a figyelmet a háború politikai természetére: Clausewitz klasszikus tétele szerint „a háború a politika folytatása más eszközökkel” [30]. A stratégiai célok és a katonai eszközök összekapcsolásának kényszere a későbbi doktrínákban is meghatározó maradt, és közvetlenül hat a tervezési folyamatok felépítésére (célok → feladatok → erők és eszközök → kockázatok).

3.3. Ipari korszak és világháborúk: a modern hadművelleti művészet kialakulása

Az ipari korszak technológiai és társadalmi változásai (vasút, távíró, tömeges mozgósítás, tömegtermelés) a hadműveletek méretét és komplexitását ugrásszerűen növelték. A vezérkari rendszerek (különösen a porosz–német hagyomány) és a részletes mozgósítási/vasúti tervek a „tervezési intézményesülés” alapját adták. A U.S. Army Center of Military History tanulmánya kiemeli, hogy a modern operatív szint nyugati-európai gyökerű: Napoleon adaptációi nyomán jelent meg a stratégiai célok és a harcászati célok közötti „köztes tér”, amelyet később a német gondolkodás (Moltke) formált elsőként koherens operatív koncepcióvá [29].

Az I. világháború állóháborús környezete a részletekbe menő, tűz- és logisztika-központú tervezést erősítette; ugyanakkor a stratégiai és operatív célok összekapcsolása gyakran nehezen volt fenntartható. A II. világháborúban az operatív művészet több irányzatban érett be (manőver, mélységi műveletek, koalíciós hadműveletek), ami a nagy hadszínterek kiterjedt, többfázisú kampányainak tervezését tette szükségessé. A stratégiai–operatív–harcászati kapcsolat gyakorlati kezelése ekkor vált a modern tervezési kultúra egyik meghatározó tapasztalatává [29].

3.4. Hidegháború és poszthidegháború: joint, expeditionary, hálózatalapú műveletek

A hidegháborúban a nukleáris elrettentés, a gyors reagálás és a NATO-szintű integrált védelem igénye a tervezés standardizálását, a kompatibilis eljárásokat és a többnemzeti parancsnoki láncok összehangolását helyezte előtérbe. A poszthidegháborús időszakban a válságkezelő és expedíciós műveletek tervezése vált dominánssá: a művelti környezet összetettebbé (állami és nem állami szereplők), a célrendszer pedig gyakran „többcélúvá” (biztonság, stabilizáció, intézményépítés) vált. A NATO tervezési doktrínák ennek megfelelően a művelti koncepció, a döntő feltételek és az értékelési logika megerősítésével reagáltak; a COPD fejlődését elemző hazai szakirodalom szerint a doktrinális háttérben több korrekció is történt, különösen a célok–hatások–döntő feltételek gyakorlati alkalmazhatóságának erősítése érdekében [31].

Az összhaderőnemi tervezés és a többnemzeti integráció doktrinális rögzítése mind a NATO, mind az amerikai doktrínában hangsúlyos. A NATO AJP-5 a művelti tervezést a katonai-stratégiai szinttől az operatív szintig, és annak harcászati szintű kapcsolódásáig írja le [27]. A

US JP 5-0 külön alfejezetben tárgyalja az operatív szintű integrációt többnemzeti műveletekben, és rámutat, hogy a parancsnokok rendszerint két parancsnoki láncban működnek (nemzeti és többnemzeti) [9].

3.5. Multi-domain és információs kor: komplex adaptív műveleti környezet

A 21. század műveleti környezetét a gyors információáramlás, az érzékelők–fegyverrendszerek–döntéshozatal összekapcsolódása, valamint a fizikai és nem-fizikai (kognitív, információs, kiber) térben zajló hatásmechanizmusok együttesen alakítják. Az amerikai tervezési doktrína a stratégiai környezetet kifejezetten „transzregionális, all-domain (szárazföldi, légi, tengeri, űr- és kiber), és multifunkcionális” kihívásokkal jellemzi [32]. Ugyanakkor a tervezésben megjelenő rendszerelvű gondolkodás és az összetettség kezelése nem csak technológiai, hanem módszertani kérdés: a JP5-0 kiemeli, hogy az operatív tervezés/operational design a műveleti környezetet komplex, kölcsönhatásokkal teli rendszerként szervezi és értelmezi, és ezzel támogatja a döntéshozatalt [9].

A multi-domain megközelítés tervezési szempontból azt jelenti, hogy a célok eléréséhez szükséges hatásokat több tartományban és több műveleti funkcióban kell szinkronizálni, miközben a szövetségesi és partnerségi együttműködés (adatmegosztás, interoperabilitás, jogi és politikai korlátok) tovább növeli a komplexitást. Ebben a környezetben a gyors helyzetértékelés, az alternatív cselekvési változatok (COA) megalapozása és a kockázatok dinamikus újraértékelése a hadműveleti tervezés kulcsképességévé válik – és közvetlen kapcsolódási pontot teremt az értekezés központi témájához, az MI-alapú tervezéstámogatáshoz.

3.6. Összefoglalás

A művelettervezés fejlődése az ókori, alapelvekre épülő hadtudománytól a napóleoni korszak manőver-központú hadjáratain át az ipari korszak és a világháborúk tömeges, többfázisú hadműveleteiig vezetett. A hidegháború és az azt követő válságkezelő korszak a standardizált, összhaderőnemi és többnemzeti tervezési eljárások megerősödését hozta, miközben az információs kor komplex, több tartományt érintő kihívásai új módszertani és technológiai igényeket támasztanak. E történeti–elméleti áttekintés alapján a következő fejezetek már

célzottan vizsgálhatják, hogy a NATO COPD módszertan egyes lépéseihez milyen MI-alapú képességek illeszthetők, és milyen korlátok mellett.

4. NATO ÉS EU MŰVELETTERVEZÉSI DOKTRÍNÁK ÉS MÓDSZERTANOK

A fejezet célja a NATO és az Európai Unió (EU) művelettervezési doktrínáinak és módszertanainak összehasonlító bemutatása, különös tekintettel a NATO-domináns gyakorlati környezetre (szövetségi parancsnoki lánc, interoperabilitási követelmények, szabványosított tervezési termékek). A fejezet (1) röviden áttekinti az átfogó megközelítés (comprehensive approach) és a DIME/PMESII keretek szerepét a katonai tervezésben; (2) ismerteti a NATO AJP-5 tervezési doktrínáját; (3) bemutatja a NATO Comprehensive Operational Planning Directive (COPD) v3.1 irányelv helyét és szerepét az AJP-5-höz viszonyítva; (4) összefoglalja az EU katonai tervezés CONOPS–OPLAN logikáját és az MPCC/EUMS szerepeit; (5) rögzíti a magyar katonai doktrínák és szabályzók fő kapcsolódási pontjait. A fejezet következtetései később alapot adnak ahhoz, hogy a COPD-lépésekhez milyen MI-alapú támogató képességek illeszthetők.

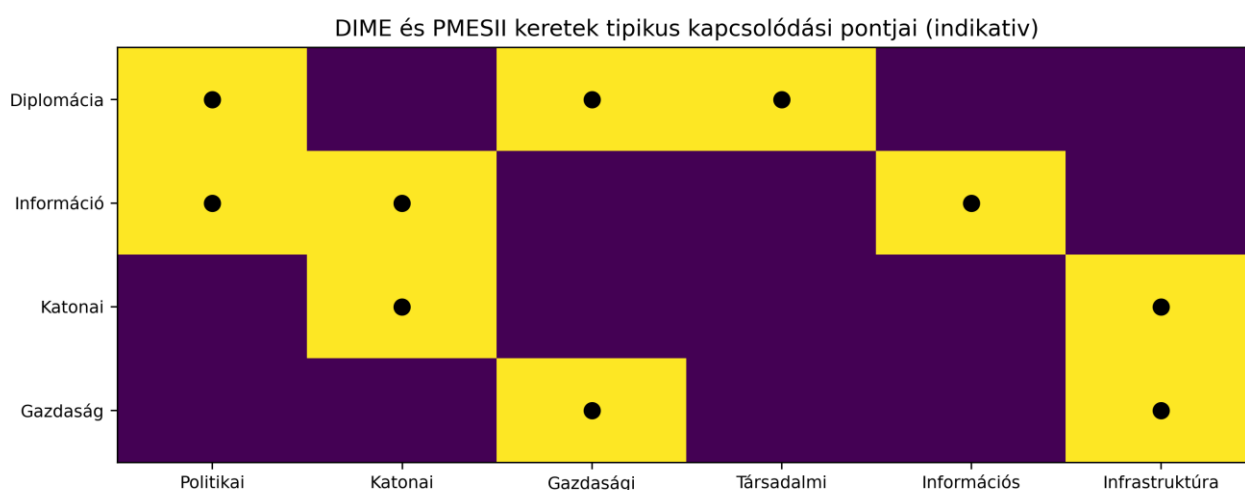
4.1. Átfogó megközelítés (comprehensive approach) és a DIME/PMESII keretrendszerek

A kortárs válságkezelési és elrettentési környezetben a katonai erő alkalmazása ritkán önmagában döntő; a politikai célok elérése több szektor együttes, időben összehangolt beavatkozását igényli. A NATO a comprehensive approach fogalmát a szövetségi tevékenységek (politikai, katonai, civil, partnerek és nemzetközi szervezetek) összehangolásának elveként kezeli, és a stratégiai dokumentumok a reziliencia és a társadalmi ellenálló-képesség erősítését a kollektív védelem és válságkezelés alapfeltételeként rögzítik [95]. AJP-01 külön fejezetben tárgyalja a NATO hozzájárulását a comprehensive approach-hoz, kiemelve a széleskörű partnernműködés és a civil–katonai koordináció szerepét [96, PDF pp. 66–70].

Az átfogó megközelítés a tervezésben azt jelenti, hogy a parancsnok és a törzs a műveleti környezetet komplex rendszerként értelmezi, és a katonai tevékenységet a rendelkezésre álló nemzeti hatalmi eszközök és partnerek tevékenységeivel (diplomácia, információ, gazdaság, belbiztonság, humanitárius támogatás) összhangban tervezi. Az AJP-5 a comprehensive approach-ot a tervezési megfontolások között rögzíti, és a műveleti környezet megértését (understanding the operational environment) a tervezés egyik kulcs-előfeltételeként kezeli [8,

PDF pp. 13, 32–34]. A COPD v3.1 ugyanezt operacionalizálja: a tervezési lépések és briefíngek beépítik a civil tényezők és partnerek elemzését, különösen a mission analysis és a műveleti megközelítés kialakításának szakaszában [30, PDF pp. 6–11, 24–29].

Az átfogó megközelítés tervezési „nyelve” gyakran két, egymást kiegészítő keretrendszeren keresztül jelenik meg. A DIME (Diplomatic, Informational, Military, Economic) a nemzeti hatalom eszközeit rendszerezi a stratégiai–politikai gondolkodás számára, míg a PMESII (Political, Military, Economic, Social, Information, Infrastructure) a műveleti környezet rendszerszintű, elemző felbontására alkalmas a felderítés–hírszerzés és a műveleti tervezés támogatásában. AJP-5 és a COPD is használja a PMESII logikát a környezet „dimenzióinak” strukturálására [8, PDF pp. 29, 57], [30, PDF pp. 28, 38–39, 45–46]. Az amerikai JIPOE-doktrína a PMESII dimenziókat a rendszerelemzés és az ellenség/ellenelő viselkedési mintázatainak feltárásához javasolja alkalmazni [10, PDF pp. 13–15].



4-1. ábra – DIME és PMESII keretek tipikus kapcsolódásai (saját szerkesztés; AJP-5 és COPD fogalomhasználat alapján [8], [30]).

Keret	Fő cél	Tipikus felhasználás a tervezésben
DIME	Nemzeti hatalmi eszközök rendszerezése	Stratégiai–politikai célok és eszközök illesztése; CA/IA integráció
PMESII	Műveleti környezet dimenzionális felbontása	OE megértése, JIPOE, hatásmechanizmusok és sérülékenységek feltárása

Comprehensive / Integrated approach	Civil–katonai és többügynökségi összehangolás	Katonai tevékenység beágyazása a szélesebb eszközrendszerbe
-------------------------------------	---	---

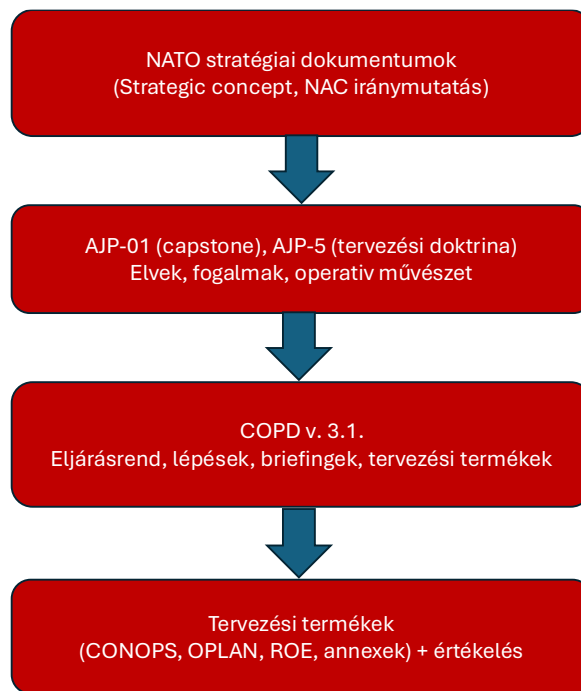
4-1. táblázat – DIME/PMESII és átfogó megközelítés: fogalmi szerepek (saját szerkesztés).

A comprehensive approach gyakorlati következménye, hogy a tervezés során a katonai célok (military end state) és hatások (effects) akkor tekinthetők megalapozottnak, ha explicit módon kapcsolódnak a politikai célokhoz, és figyelembe veszik a nem katonai korlátokat (jogi, politikai, társadalmi) és a partnerek tevékenységét. A NATO 2022-es Stratégiai Koncepciója a reziliencia és a whole-of-society megközelítés fontosságát külön is kiemeli, ami a tervezésben a kritikus infrastruktúrák és társadalmi alrendszerek védelmének prioritizálását támasztja alá [95, pp. 4–5]. Az EU hasonló logikát „integrated approach” néven rögzít, hangsúlyozva a konfliktusok és válságok kezelésének holisztikus, többeszközű keretezését [97, PDF pp. 1–4].

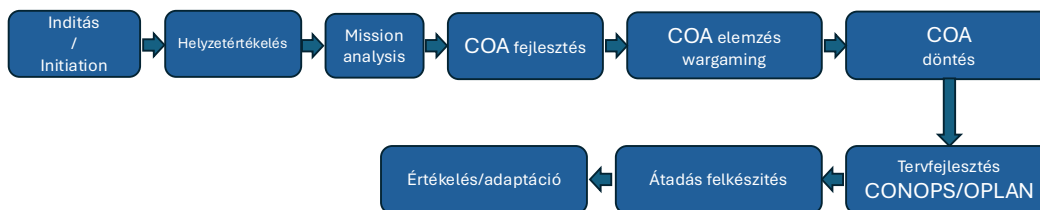
4.2. NATO AJP-5: a hadművelettervezés doktrínája

Az AJP-5 a NATO művelettervezésre vonatkozó autoritativ doktrínája: meghatározza a tervezés alapelveit, fogalmait és a tervezési logika doktrinális kereteit. Szerepe kettős. Egyrészt az AJP-01 capstone doktrínához illeszkedve a stratégiai–hadműveleti–harcászati szintek közötti kapcsolatot, az operatív művészetet és a műveletek keretezését tárgyalja [96, PDF pp. 32–38]. Másrészt a konkrét tervezési eljárásokhoz (OPP, tervtermékek, wargaming, kockázatkezelés) doktrinális háttérrel ad, amelyet a COPD v3.1 részletes eljárási irányelveként bont ki [8], [30].

Az AJP-5 a tervezést iteratív, bizonytalanságban zajló döntéstámogató folyamatként kezeli. A doktrína hangsúlyozza, hogy a parancsnoki „design” (konceptcionális keretezés) és a részletes tervekészítés egymást kiegészítő tevékenységek: a design biztosítja a probléma helyes értelmezését és a műveleti megközelítés (operational approach) koherenciáját, míg a részletes tervezés a végrehajtásra alkalmas OPLAN/annex struktúrát hozza létre [8, PDF pp. 23–26, 57–61]. E logika a JP 5-0 amerikai összhaderőnemi doktrínával is összhangban van, amely az operational design-t a komplex műveleti környezet rendszerszintű értelmezésének módszertanaként írja le [7, PDF pp. 5–14].



4-2. ábra – NATO tervezési „architektúra”: stratégiai dokumentumok → AJP-5 doktrína → COPD irányelv → termékek (saját szerkesztés [8], [30] alapján).



4-3. ábra – NATO OPP ciklus vázlata (saját szerkesztés; AJP-5 és COPD alapján [8], [30]).

Az AJP-5 kulcsfogalmi készlete a műveleti tervezésben néhány visszatérő elem köré szerveződik. A „centre of gravity” (COG) a döntő képesség/forrás vagy rendszer, amelynek sérülékenységei révén az ellenség működése vagy a saját képességek megvédhetők; a doktrína

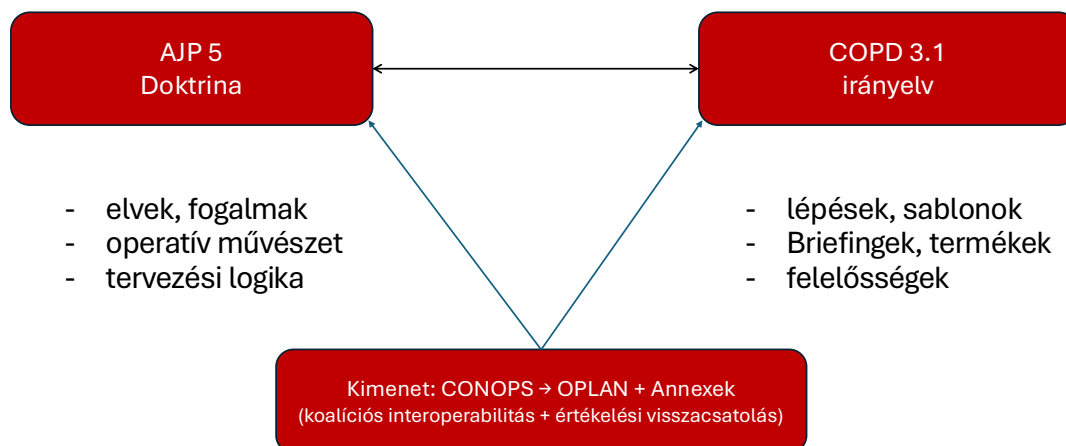
kiemeli, hogy a COG-elemzés a környezet megértéséhez és a döntő feltételek (decisive conditions) azonosításához járul hozzá [8, PDF pp. 56–61]. A decisive conditions a kívánt végállapot eléréséhez szükséges, időben és térben elérendő feltételek rendszere; ezek a LOE/LOO logikával összekapcsolva adják a műveleti megközelítés gerincét [8, PDF pp. 53–66].

A lines of operation (LOO) és lines of effort (LOE) fogalompár a műveleti tevékenységek szervezésére szolgál. A LOO tipikusan tér-idő jellegű, „földrajzi” logika mentén szervezi a műveleteket, míg a LOE inkább tematikus, hatásorientált (pl. kormányzás, biztonság, gazdasági stabilizáció) szervező elv, amely különösen a comprehensive approach környezetében hasznos [8, PDF pp. 53–66]. AJP-5 hangsúlyozza, hogy a tervezési termékeknek a stratégiai célokhoz való visszavezethetősége (traceability) és a kockázatok explicit kezelése kritikus: a COA-k összevetésekor a kockázat a döntés integrált része, nem utólagos megfontolás [8, PDF pp. 23–26, 44–47].

A tervezésben az „assessment” (értékelés) az adaptív ciklus biztosítója: a végrehajtás során a parancsnoknak visszajelzést kell kapnia arról, hogy a tevékenységek a kívánt hatások felé visznek-e, és szükséges-e a terv módosítása. Az AJP-5 ezt a tervezési logikába integrálja, és a NATO Operations Assessment Handbook részletesen kifejti a MOP/MOE (Measures of Performance/Effectiveness) és a visszacsatolási mechanizmusok alkalmazását a COPD-vel összefüggésben [98, PDF pp. 15–18, 35–40]. A doktrinális következmény: a tervezési termékek (CONOPS/OPLAN és annexek) már a tervezés során olyan indikátorokat és adatforrásokat kell, hogy kijelöljenek, amelyek később mérhetővé teszik a műveleti előrehaladást.

4.3. NATO COPD v3.1: cél, helye és kapcsolódása az AJP-5-höz

A COPD v3.1 a NATO művelettervezés részletes irányelve (directive), amely az AJP-5 doktrinális elveit és fogalmait konkrét eljárásokra, tervezési termékekre, briefingekre és felelőségekre bontja le. Míg az AJP-5 „mit és miért” kérdésekre ad doktrinális választ (fogalmi készlet, operatív művészet, tervezési alapelvek), addig a COPD „hogyan” jelleggel standardizálja a tervezési folyamat lépéseit, támogatva a többnemzeti törzsek interoperabilitását [8], [30].



4-4. ábra – AJP-5 (doktrína) és COPD (irányelv) funkcionális kapcsolata (saját szerkesztés [8], [30] alapján).

A COPD a NATO Crisis Response System (NCRS) és a NATO Crisis Response Process (NCRP) logikájába ágyazva írja le a stratégiai és operatív szint közötti tervezési kapcsolódást. A dokumentum a tervezési fázisokat és a termékek láncolatát úgy kezeli, hogy a NAC/MC szintű politikai–katonai iránymutatásból levezethető legyen a stratégiai CONOPS/OPLAN, majd a hadművelleti szintű OPLAN és részletes annex rendszer [30, PDF pp. 10–18]. A döntési pontok és briefingek (pl. mission analysis briefing, COA briefing, plan review) a többnemzeti törzs együttműködésének „ritmusát” adják, és a parancsnoki döntéstámogatás standardizált csatornái [30, PDF pp. 11–18, 29].

A COPD v3.1 erőssége, hogy a koncepcionális tervezés (operational design) és a részletes tervezési termékek közötti átjárást konkrét sablonokkal és követelményekkel támogatja. A mission analysis szakaszban például a környezet megértése (PMESII dimenziók), a COG-elemzés, a decisive conditions és a kockázatok beazonosítása összekapcsolódik a parancsnoki szándék és a műveleti megközelítés kialakításával [30, PDF pp. 24–29, 38–39, 49]. A DIME/PMESII keretek alkalmazása a COPD-ben nem „öncélú elemzés”, hanem a döntéselőkészítés strukturálása: a tervezői közösségnek közös nyelvet ad a nem katonai tényezők és a katonai tevékenységek kapcsolatának megragadására [30, PDF pp. 38–39, 45–46].

A COPD kiemeli az értékelési visszacsatolást (assessment) mint a tervezés és végrehajtás összekötő elemét. A dokumentum több ponton utal az operations assessment szerepére, és a NATO Operations Assessment Handbook kifejezetten a COPD-vel együtt értelmezendőként fogalmazza meg az értékelési eljárásokat [98], [30, PDF pp. 11, 16–18]. A MOP/MOE logika,

a baseline meghatározása, valamint az indikátorok kijelölése (adatforrások, reporting) a tervtermékek integráns része kell legyen, különösen komplex, hosszú távú stabilizációs vagy elrettentési műveletekben.

Szempon	AJP-5 (doktrína)	COPD v3.1 (irányelv)
Fő szerep	Elvi–fogalmi keret, operatív művészet, tervezési logika	Eljárásrend, lépések, briefingek, termékek
Kimenet	Tervezési alapelvek és fogalmak a szövetségi alkalmazáshoz	Standardizált CONOPS/OPLAN-termékek és annexek
Interoperabilitás	Közös doktrinális nyelv	Közös munkafolyamat és dokumentumsablonok
Assessment	Doktrinális beágyazás	Konkrét beépítés termékekbe és briefingekbe

4-2. táblázat – AJP-5 és COPD funkcionális különbségei (saját szerkesztés [8], [30] alapján).

A „doktrína–irányelv” különbség különösen fontos a disszertáció későbbi (MI-támogatás) fejezetei szempontjából. A COPD szintjén a tervezési folyamat már jól strukturált feladatokra és termékekre bontható (pl. PMESII elemzés, COG azonosítás, COA kidolgozás és wargaming, MOP/MOE kijelölés), ami technológiai támogatásra – így MI-alapú döntéstámogatásra – „illeszthetőbbé” teszi a módszertant. Ugyanakkor a doktrinális szinten (AJP-5) rögzített parancsnoki felelősség, kockázatkezelés és comprehensive approach követelmények meghatározzák a támogatás lehetséges határait és a szükséges kontrollokat.

4.4. EU katonai tervezés: CONOPS–OPLAN logika, MPCC/EUMS szerepek

Az EU közös biztonság- és védelempolitikájának (CSDP) katonai műveletei és missziói sajátos intézményi és jogi keretben zajlanak: politikai kontroll és stratégiai irányítás jellemzően a PSC-n keresztül, katonai tanácsadás és iránymutatás az EUMC-n, míg a katonai törzsmunka és a tervezési termékek előállítása az EUMS/MPCC és az adott műveleti parancsnokság (OpCdr/OHQ) feladata. Az EU „integrated approach” a katonai tervezést is a szélesebb EU-eszköztárba ágyazza, hangsúlyozva a politikai, diplomáciai, fejlesztési és humanitárius eszközök koherens alkalmazását [97, PDF pp. 1–4], [99, PDF p. 11].



Intézményi keret (tipikus) : PSC/EUMC irányítás, EUMS/MPCC tervezés és vezetés (nem executive), lásd ST-6432/2015 , ST-8798/2019

4-5. ábra – EU CSDP katonai tervezés: PFCA/CMC → CONOPS → OPLAN (saját szerkesztés; EU ST-6432/2015 alapján [99]).

Az EU katonai tervezés politikai–stratégiai szintű koncepcióját a Council dokumentumok rögzítik. Az EU Concept for Military Planning at the Political and Strategic Level (ST-6432/2015) a tervezést iteratív folyamatként írja le, amely a politikai keretektől (PFCA/CMC) indul, majd a Military Strategic Options (MSO) és a Initiating Military Directive (IMD) után a CONOPS és az OPLAN kidolgozásáig jut el [99, PDF pp. 11–12]. A dokumentum a tervezést a comprehensive approach logikájába ágyazza, és kiemeli, hogy a katonai tervezés nem választható el az EU többi instrumentumától [99, PDF p. 1].

Az EU parancsnoki és vezetési (C2) rendjét a ST-8798/2019 részletezi. A dokumentum hivatkozik a 2017/971 tanácsi határozatra, amely létrehozta az MPCC-t az EUMS keretén belül, és a nem-executive katonai missziók tervezéséért és vezetéséért a Director MPCC felelősségét rögzíti [100, PDF pp. 1–3], [101, PDF pp. 1–3]. A ST-8798/2019 több ponton részletezi az OpCdr/MCdr szerepeit, a jelentési rendet és a tervezési termékek (CONOPS/OPLAN) követelményeit [100, PDF pp. 7–12, 17–20].

A NATO-hoz viszonyítva az EU tervezési logika több szempontból párhuzamos: mindkét szervezet a CONOPS–OPLAN láncban strukturálja a végrehajtásra érett tervtermékeket, és mindkettőben erős a többszintű (politikai–katonai) kontroll. Ugyanakkor eltérés, hogy az EU CSDP műveletek jellemzően szűkebb katonai skálán és szigorúbb politikai korlátozások mellett zajlanak; a „command options” (pl. nemzeti OHQ-k, MPCC) rugalmasabbak, de a NATO-hoz képest kisebb a hosszú távon bejáratott szabványosítás mélysége. Emiatt az

interoperabilitás biztosítása (különösen NATO–EU párhuzamos szerepvállalásoknál) gyakran nem doktrinális, hanem szervezeti és információmegosztási kérdésként jelenik meg.

Szempont	NATO (AJP-5/COPD)	EU (ST-6432/2015, ST-8798/2019)	Megjegyzés
Fő termék-lánc	CONOPS → OPLAN (+ annexek)	CMC/IMD → CONOPS → OPLAN	Mindkettő iteratív, briefing-alapú
Intézményi kontroll	NAC/MC + SHAPE/JFC	PSC/EUMC + EUMS/MPCC	Politikai kontroll mindkét esetben erős
Assessment	NATO OAH + COPD integráció	CSDP reporting és értékelés	NATO-nál erősen formalizált MOP/MOE

4-3. táblázat – NATO és EU tervezési logika összevetése (saját szerkesztés [30], [99], [100] alapján).

4.5. Magyar katonai doktrinák és szabályzók kapcsolódásai

A Magyar Honvédség (MH) doktrinális és eljárásrendi környezete a NATO-hoz való csatlakozás óta a szövetségi interoperabilitási követelményekhez igazodik. A nemzeti katonai stratégia a szövetségi kötelezettségek és a kollektív védelem prioritását hangsúlyozza, miközben rögzíti a képességfejlesztés és a válságkezelési feladatok kereteit [11, PDF pp. 3–10]. A tervezés szempontjából ez azt jelenti, hogy a magyar eljárásoknak kompatibilisnek kell lenniük a NATO tervezési termékeivel (CONOPS/OPLAN, annex struktúrák) és a többnemzeti törzsmunka ritmusával (briefingek, döntési pontok).

Az MH törzsszolgálati szakutasítás több helyen hivatkozik a NATO eljárásokra és a művelettervezés standardizált termékeire; a dokumentumban a „COPD” kifejezés is explicit módon megjelenik, ami arra utal, hogy a tervezési folyamat nem pusztán „általános elvekben”, hanem konkrét NATO irányelvi szinten is beépül a nemzeti törzsmunka rendjébe [54, PDF pp. 7, 60, 154–160]. Ez különösen a koalíciós környezetben releváns, ahol a magyar törzseknek a NATO-ritmusban kell tervezési termékeket előállítaniuk és értékelniük.

Az Összhaderőnemi Doktrína (Ált/44) célja, hogy az MH műveleteinek alapelveit a szövetségi doktrinákhoz igazítsa; a dokumentum több helyen hivatkozik NATO kiadványokra (AJP sorozat), és a joint szemléletet, valamint az összhaderőnemi funkciók integrációját a műveletek

kulcselemeként kezeli [102, PDF pp. 7–12, 62]. A nemzeti szabályzók gyakorlati kihívása ugyanakkor, hogy a NATO doktrínák és irányelvek (pl. AJP-5, COPD) időről időre frissülnek; ezért a nemzeti dokumentumok karbantartása és a törzsek képzése (terminológia, terméksablonok, szoftveres támogatás) folyamatos feladat.

A magyar kapcsolódások értékelésénél három dimenzió különösen fontos a disszertáció MI-fókuszra szempontjából. Elsőként: az interoperabilitási kényszer a NATO termékek és folyamatok „gépesíthető” elemeinek átvételét támogatja (sablonok, adatstruktúrák, termékklánc). Másodszor: a nemzeti korlátozások és minősítési szintek adatmegosztási és információbiztonsági korlátokat teremtenek, amelyek később az MI-rendszerek bevezetését is korlátozzák (adat-hozzáférés, audit, traceability). Harmadszor: az MH-ban a törzsek terhelése és a szűk humán erőforrás a döntéstámogató automatizálás irányába mutat – de csak a NATO doktrínákkal konzisztens governance és minőségbiztosítás mellett.

4.6. Összefoglalás

A fejezet bemutatta, hogy a NATO és az EU művelettervezési rendszereinek közös nevezője a politikai célok és a katonai eszközök összehangolása, valamint a többszereplős válságkezelési környezet átfogó (comprehensive/integrated) keretezése. A NATO esetében az AJP-5 doktrína biztosítja a fogalmi és elvi alapokat (operatív művészet, design, COG, decisive conditions, LOE/LOO), míg a COPD v3.1 a standardizált folyamatlépések és termékek révén teszi végrehajthatóvá és interoperábilissá a többnemzeti törzsmunka gyakorlatát [3], [4]. Az EU tervezés CONOPS–OPLAN logikája szerkezetében hasonló, de intézményi kerete és műveleti skálája eltérő; az MPCC létrehozása a katonai C2 és tervezés központosításának fontos lépése [99]–[101].

A magyar doktrínák és törzsszolgálati szabályzók a NATO kompatibilitást alapelveként kezelik, és a nemzeti tervezési gyakorlatban a NATO eljárások (így a COPD) használata kimutatható [54], [102]. A fejezet következtetése, hogy az MI-alapú tervezéstámogatás szempontjából a NATO-domináns környezet különösen kedvező: a COPD lépései és termékei világosan strukturáltak, így a későbbi fejezetekben megalapozottan vizsgálható, hogy mely tervezési részfeladatoknál (információgyűjtés, OE-elemzés, COG/decisive conditions támogatás, COA-értékelés, assessment indikátorok) alkalmazható MI – és milyen doktrinális, jogi és biztonsági korlátok mellett.

5. A NATO COMPREHENSIVE OPERATIONAL PLANNING DIRECTIVE (COPD) V3.1 MÓDSZERTAN RÉSZLETES ELEMZÉSE

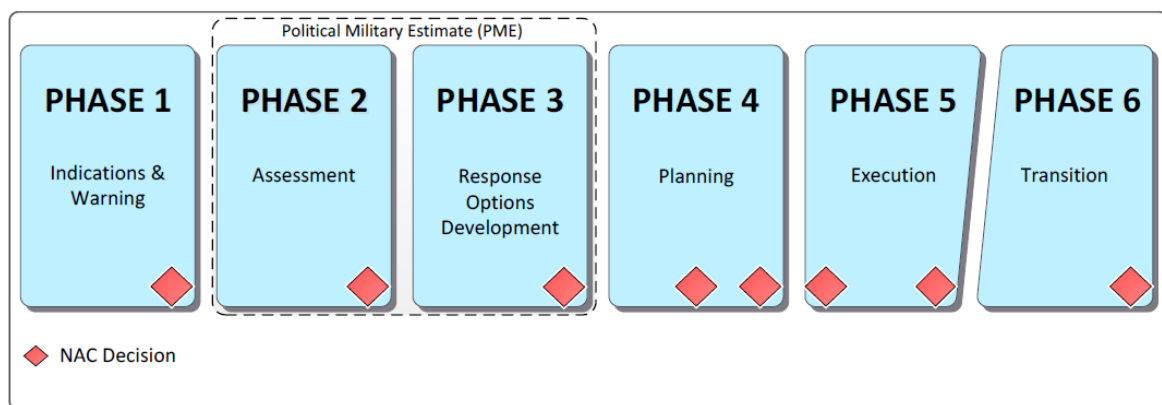
Ez a fejezet a NATO Szövetséges Műveleti Parancsnokság (Allied Command Operations – ACO) művelettervezési keretrendszerének egyik kulcsdokumentumát, a Comprehensive Operational Planning Directive (COPD) v3.1-et elemzi. A COPD az AJP-5 (Allied Joint Doctrine for the Planning of Operations) doktrinális elveire épít, ugyanakkor annál részletesebb és „működés-közelí”: a hadműveleti szintű törzsek mindennapi tervezési munkáját szabályozza. A dokumentum egységes fázislogikát, standardizált termékstruktúrát, ellenőrző pontokat (decision points) és döntéstámogató eszközöket (pl. Decision Support Template/Matrix) határoz meg. Célja, hogy a tervezés a többnemzeti környezetben is konzisztens, auditálható és végrehajtható legyen, miközben a parancsnoki döntés-előkészítést helyezi a középpontba [33].

A COPD szemléletének középpontjában a commander-centric planning áll: a parancsnok nem a folyamat végén „átveszi” a tervet, hanem a tervezés egészét irányítja a megfelelő időben kiadott iránymutatásokkal. A törzs feladata nem az, hogy a lehető legrészletesebb dokumentumot állítsa elő, hanem hogy a döntési helyzetet világossá tegye: milyen opciók állnak rendelkezésre, ezek milyen kockázatokkal járnak, és milyen következmények várhatók a műveleti környezetben. A COPD ezért a tervezési termékeket döntési csomagokként kezeli, amelyeknek visszavezethetőnek kell lenniük a stratégiai célokra, és amelyekben a feltételezések és kockázatok explicit módon szerepelnek [33], [34].

A fejezet célja, hogy a COPD v3.1 folyamatát és termékeit olyan részletezettséggel bontsa ki, amely később a disszertáció MI-támogatási fejezeteiben konkrét beavatkozási pontokhoz kapcsolható. Az elemzés a tervezési munka kritikus szakaszaira fókuszál: (1) helyzetértés és probléma-keretezés (ISA/SSA), ahol adat- és szintézishiányok keletkezhetnek; (2) COA-fejlesztés és wargaming, ahol kognitív torzítások és csoportdinamikai problémák jelentkezhettek; (3) döntéstámogatás és assessment, ahol az indikátorok kiválasztása és az információátterhelés okoz kockázatot; valamint (4) lezárás és lessons learned, ahol a szervezeti tanulás intézményesítése történik [33], [35], [36].

A fejezet szerkezete a COPD logikáját követi. Az 5.1 alfejezet ismerteti az alapelveket, a szerepköröket és a tervezési termékek rendszerét (SSA–MRO–CONOPS–OPLAN, döntési

pontok, traceability). Az 5.2 alfejezet fázisonként (Phase 1–6) elemzi a tevékenységeket, döntési csomagokat és a gyorsított (accelerated) tervezés kompromisszumait. Az 5.3 alfejezet az operational design módszertani magját tárgyalja (COG, DC, LoE/LoO, hatásláncok, kockázat). Az 5.4 a COA-fejlesztés, wargaming és összehasonlító értékelés módszerét bontja ki, az 5.5 az assessment és lessons learned beágyazását elemzi, az 5.6 pedig a jelen kihívásokat foglalja össze. Végül az 5.7 összegzi a fejezet legfontosabb következtetéseit és az MI-támogatás szempontjából releváns pontokat.



ISA

SSA

MRO

CONOPS/OPLAN

EXEC

Ábra 5-1. A COPD tervezési fázisai és fő termékei, döntési pontok logikájával (COPD v3.1 alapján) [33].

5.1. COPD alapelvek, szerepkörök, tervezési termékek rendszere

A COPD egyik központi alapelve a parancsnoki döntés-előkészítés elsődlegessége. A tervezés célja nem az, hogy „teljes” dokumentációt készítsen, hanem hogy a parancsnok számára a releváns alternatívák és kockázatok átláthatóvá váljanak. Ebből következik a fázislogika és a döntési pontok rendszere: minden fázis végén olyan döntési csomag készül, amely rögzíti a problémakeretet, a feltételezéseket, az opciókat és a kockázatvállalást. A döntési pontok ugyanakkor minőségbiztosítási kapuk is: csak akkor érdemes a részletekbe lépni, ha a koncepció és a logikai lánc (cél–hatás–feladat) koherens. A COPD ezzel csökkenti annak kockázatát, hogy a törzs a részletezésben „elrejt” a stratégiai ellentmondásokat [33], [34].

A második alapelv a shared understanding, azaz a közös helyzetértés és közös fogalmi készlet megteremtése. Többnemzeti törzsekben a félreértések tipikusan nem a végcélok szintjén, hanem a köztes fogalmak (objective, decisive condition, effect, success criteria) eltérő értelmezésében keletkeznek. A COPD ezért erős standardizációt alkalmaz: sablonok, minimum-tartalmi elvárások és terminológiai fegyelem révén igyekszik biztosítani, hogy a funkcionális területek (J2–J9) ugyanazt a „dokumentum-nyelvet” beszéljék. A sablonok nem a kreativitás korlátozását szolgálják, hanem azt, hogy a döntési briefek összehasonlíthatók legyenek, és az integrált tervezés (integrated staff work) ne csupán egymás mellé helyezett szakági anyagokból álljon [33].

5.1.1. Tervezési szervezet: JOPG, working groupok és boardok.

A COPD a Joint Operations Planning Group (JOPG) köré szervezi a tervezést, mint integráló magcsoportot. A JOPG-ban a tervezés vezetése mellett jelen kell lennie a hírszerzési (J2), műveleti (J3), logisztikai (J4), kommunikációs és informatikai (J6), jogi (Legal), információs (STRATCOM/Info Ops), valamint civil dimenziót képviselő (J9/CIMIC) szakértőknek. A JOPG feladata a termékek koherenciája: biztosítja, hogy az SSA-ban rögzített feltételezések és korlátok megjelenjenek a COA-kban és az OPLAN-ban; illetve hogy az annexek ne mondjanak ellent a fő tervlogikának. A JOPG mellett tematikus working groupok működnek (pl. fires/targeting, information, cyber, logistics), amelyek a szakági mélymunkát végzik; a boardok pedig a döntések előkészítését és a battle rhythm szerinti jóváhagyásokat szolgálják. A szervezeti modell célja a párhuzamosítás: a részfeladatok külön csoportokban futnak, miközben a JOPG integrálja az eredményeket [33], [34].

5.1.2. Battle rhythm, „quality gates” és a tervezési minimum fogalma.

A battle rhythm az a szervezési mechanizmus, amely biztosítja, hogy a tervezés ütemezetten és ellenőrzöten haladjon. A COPD fázisok végén minőségi kapukat alkalmaz: például a mission analysis brief jóváhagyása nélkül nem indokolt COA-k kidolgozásába kezdeni, mert a probléma és end state még nem stabil. A quality gate lényege nem formai, hanem tartalmi: a döntéshez szükséges minimum-információ és döntési logika legyen jelen (probléma, opciók, feltételezések, kockázat). Gyorsított tervezésben a minőségi kapuk akár összevonhatók, de a döntési minimum elemei nem hagyhatók el: ha a feltételezések és

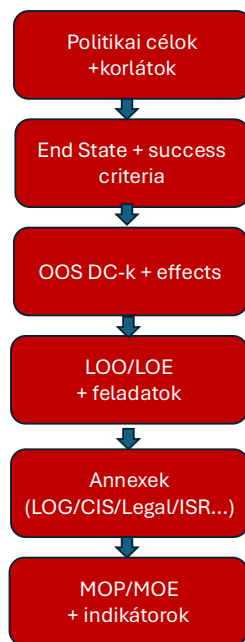
kockázatok nincsenek dokumentálva, a végrehajtás alatti adaptáció (FRAGO) kontrollálatlanná válik, és az elszámoltathatóság (accountability) is sérül [33].

5.1.3. Tervezési termékek rendszere: SSA–MRO–CONOPS–OPLAN.

A COPD termékrendszere egymásra épülő döntéstámogató dokumentumokból áll. Az SSA (Strategic Assessment) rögzíti a műveleti környezet lényegi elemeit, az end state-et és success criteria-t, valamint a korlátozásokat és feltételezéseket. Az MRO (Military Response Options) a lehetséges katonai válaszopciók strukturált összevetése: erő- és időigény, kockázat, jogi/politikai elfogadhatóság és várható következmények szerint. A CONOPS a kiválasztott opció operational design-ját formalizálja (COG/DC/LoE/LoO, phasing, main effort, kezdeti assessment). Az OPLAN/OPORD pedig végrehajtható tervvé alakítja a koncepciót: feladathozzárendelés, ütemezés, szinkronizáció, erő- és képességrendelés, valamint a támogató annexek integrációja. A termékek közötti kapcsolatot a traceability elv biztosítja: az OPLAN-ban szereplő feladatoknak vissza kell vezethetőknek lenniük a CONOPS DC-jeihez és az SSA end state-jéhez [33], [34].

5.1.4. Traceability és auditálhatóság.

A traceability azt jelenti, hogy a terv bármely eleme (feladat, ütemezés, erőforrás döntés) visszavezethető a magasabb szintű célokra, és fordítva: a célokhoz tartozó hatások és feladatok is azonosíthatók. A traceability több szinten hasznos. Minőségbiztosítás, mert segít kiszűrni a „feladat a feladatért” jellegű elemeket. Kommunikációs eszköz, mert a parancsnok és a politikai vezetés számára világossá teszi, hogyan szolgálják a műveleti lépések a stratégiai célt. Végül szervezeti tanulás alapja, mert a döntési naplóval együtt lehetővé teszi a későbbi értékelést: mely feltételezések bizonyultak tévesnek, hol voltak a kockázati vakfoltok, és mely döntési pontoknál kellett volna más trigger alkalmazni. A COPD ezért a traceability-t nem opcionális „szépségkritériumnak”, hanem a többnemzeti tervezés működőképességének feltételének tekinti [33], [36].



Ábra 5-2. Traceability: célok–design–feladatok–annexek–mérőszámok visszavezethetőségének logikája (saját szerkesztés) [33], [34].

A terv végrehajthatóságát a base plan mellett elsősorban az annexek és függelékek biztosítják. Ezek tartalmazzák a részletes szakági koncepciókat (fenntartás, CIS, legal, információs tevékenységek, cyber stb.), amelyek nélkül a terv sokszor csak szándéknyilatkozat marad. A többnemzeti környezetben az annexek adják azt a szerződés-szerű keretet, amelyben a nemzeti hozzájárulások összehangolhatók és ellenőrizhetők.

Annex-terület (példa)	Kulcskérdés tervezéskor	Tipikus kapcsolódás
ISR / Intelligence support	Mit kell tudni a döntéshez? Melyek a kritikus hiányok?	gap/CCIR; SSA; assessment indikátorok
LOG / Sustainment	Fenntartható-e a tempo és az operational reach? Hol a kulminacio?	MRO/CONOPS erő- és időigény; OPLAN fenntartási terv
CIS / C2	Interoperábilis és védett-e az információáramlás?	CONOPS előfeltételek; OPLAN végrehajtási feltételek
Legal / ROE	Mi a jogi keret és a korlátozások?	MRO elfogadhatóság; targeting/detenció
Force Protection	Mekkora a veszteségtűrés és a védelem követelménye?	risk register; branch/sequel

Information / STRATCOM	Mi a narratíva, és hogyan mérjük a hatást?	LoE; MOE-k; reputációs kockázat
Cyber	Milyen védelmi/támadó intézkedések szükségesek?	C2 függőségek; sebezhetőségek
Medical	Milyen egészségügyi kapacitás és MEDEVAC kell?	tempo; veszteség-modellek

5-1. táblázat. Annex-taxonómia és tipikus tervezési kapcsolódások (saját szerkesztés) [33], [34].

5.1.5. Információmegosztás, besorolás és „releasability” rétegezés.

A többnemzeti tervezés egyik legnagyobb fékező tényezője az információmegosztás korlátozottsága. A hírszerzési anyagok, különösen a forrásvédett (source-sensitive) információk gyakran nem adhatók át a teljes törzs számára, ugyanakkor a tervezéshez közös helyzetértés kell. A COPD gyakorlati megoldásként a termékek rétegezését támogatja: ugyanazon döntési csomagnak legyen egy széles körben megosztható (releasable) verziója, amely a döntési logikát és következtetéseket tartalmazza, és legyen külön kezelve a bizonyítékok szenzitív része. Ez a megközelítés csökkenti a „black box” érzést a koalíciós partnerekben, miközben tiszteletben tartja a nemzeti caveat korlátokat. A tervezési gyakorlatban ez különösen az SSA és az assessment területén kritikus, mert ott sokszor heterogén forrásokból kell közös narratívát felépíteni [33].

5.1.6. Tudásmenedzsment, verziókezelés és változásnapló.

A COPD implicit módon feltételezi, hogy a tervezési termékek élő dokumentumok: az SSA frissül az OE változásával, a COA-k a wargaming után módosulnak, az OPLAN pedig a végrehajtás során FRAGO-kkal iterál. Ez csak akkor kezelhető, ha a törzs dokumentumfegyelmet tart fenn: baseline kijelölés, változásnapló vezetése, felelős szerkesztők kijelölése, valamint kontrollált terjesztés (distribution). A leggyakoribb probléma a párhuzamos verziók keletkezése, amikor különböző working groupok eltérő verziókban dolgoznak, vagy a briefing anyag nincs szinkronban az annexekkel. A COPD szemléletében a verziófegyelem nem adminisztratív teher, hanem a döntési transzparencia és a végrehajthatóság feltétele [33].

5.1.7. Döntési napló és feltételezésmenedzsment.

Időnyomás alatt a törzs hajlamos „implicit” feltételezésekre: olyan állításokra, amelyekről mindenki azt hiszi, hogy igaz, de valójában nem ellenőrzött. A COPD logikája szerint a feltételezéseket (assumptions) explicit listába kell gyűjteni, felelőst (owner) rendelni hozzájuk, és validációs tervet készíteni (mikori adat, milyen forrás, melyik board/jelentés igazolja vagy cáfolja). A döntési napló hasonló fegyelmet hoz: rögzíti a döntési pontot, az alternatívákat, a választott opció indokát és a kapcsolódó kockázatokat. Ez a fegyelem segít a Phase 5 adaptációban is, mert láthatóvá teszi, mely döntések milyen feltételezésekre épültek; ha a feltételezés megdől, a terv mely részeit kell először újragondolni. E nélkül a FRAGO-k „reaktív foltozássá” válnak, és a terv koherenciája szétesik [33], [36].

5.1.8. Terminológiai kontroll és fogalmi konzisztencia. A többnemzeti tervezés egyik rejtett költsége a terminológiai eltérés. Ha egy nemzet az „effect” fogalmat feladatként értelmezi, míg mások állapotként, akkor az assessment és a traceability már a dokumentumszerkezet szintjén hibás lesz. A NATO Terminology Database (NATOTerm) használata és a COPD/AJP-5 szerinti fogalomhasználat a tervezés gyorsaságát is növeli: rövid távon a fogalmi tisztázás időigényesnek tűnhet, de hosszú távon csökkenti az újramunkát és a félreértésekből fakadó vitákat. A gyakorlatban célszerű a fő tervezési termékekben (SSA, CONOPS, OPLAN) egy rövid „terminology control” részt fenntartani, amely rögzíti a kulcsfogalmak alkalmazott értelmét, különösen, ha a műveletben több szervezet és partner vesz részt [37].

5.2. Fázisok és főtermékek (Phase 1–6) – döntési pontok és jóváhagyások

A COPD a tervezést fázisokra bontja, hogy a törzs munkája konvergáljon, és a tervezés során ne vesszen el a fókusz. A fázisok nem merev lépcsők: a környezet változása vagy a magasabb szint új iránymutatása miatt visszalépések és újraértékelések előfordulhatnak. A fázisok mégis kulcsfontosságúak, mert kijelölik, melyik időpontban milyen döntést kell előkészíteni, és milyen minimum-termék szükséges ehhez. A fázisok végén található döntési pontok (decision points) biztosítják, hogy a parancsnoki jóváhagyás és a kockázatvállalás dokumentált legyen [33].

A fáziselemzésnél a fejezet a következő dimenziókat használja: (1) cél és döntési kérdés; (2) tipikus inputok és szakági szerepek; (3) elemzési és szintézis tevékenységek; (4) kimeneti termékek és minőségkritériumok; valamint (5) gyakori buktatók. A checklistek a COPD-ből levezetett „tervezési minimumot” rögzítik, különösen gyorsított tervezési helyzetben. A checklistek használata nem helyettesíti a szakmai ítéletet, de csökkenti annak kockázatát, hogy időnyomás alatt a törzs kihagyjon kritikus lépéseket [33], [34], [10].

5.2.1. Phase 1 – Initial Situational Awareness (ISA) és tervezésindítás

A Phase 1 a tervezés indításának fázisa. Itt a törzs a válságtriggert (esemény, jelzés, politikai feladat) döntési kérdéssé formálja: mit kell rövid időn belül meghatározni ahhoz, hogy a tervezés érdemben elinduljon. A cél nem a teljes helyzetkép, hanem a tervezési minimum létrehozása: a művelet lehetséges kerete, a vizsgálandó és kizárt területek, a kezdeti idővonal, valamint a kritikus információhiányok (gaps) azonosítása. A COPD szerint a Phase 1 legfontosabb outputja a priorizált gap/CCIR lista és az ehhez igazított hírszerzési gyűjtési fókusz. Ha a hiánylista nem készül el, a Phase 2-ben a törzs a rendelkezésre álló adatok „foglya” lesz, és a döntés szempontjából releváns információk késve érkeznek [33].

A Phase 1-ben kezd kialakulni a battle rhythm és a tervezési munkamegosztás. A JOPG feláll, kijelöli a vezető szerkesztőket és a szakági kapcsolattartókat. A tervezés kezdetén különösen fontos a dokumentum- és adatkezelési szabályok rögzítése (mappastruktúra, verziókódolás, change log), mert a későbbi fázisokban a párhuzamos munka gyorsan verziókáoszhoz vezethet. A Phase 1-ben érdemes rögzíteni a döntési csomagok formátumát is: rövid, döntésorientált briefek, amelyek egyértelműen elválasztják a tényeket, a feltételezéseket és a szakmai értékelést. Ez a fegyelem a gyorsított tervezésben különösen fontos, mert a „később majd pontosítjuk” hozzáállás kontrollálatlan kockázatot hoz létre [33].

Ellenőrző kérdések (checklist):

- A trigger esemény és a döntési kérdés világos és rögzített?
- Meghatározták a tervezés scope-ját és a kizárt területeket?
- Felállt a JOPG, és kijelölték a fő szakági felelősöket?
- Van priorizált gap/CCIR lista felelősökkel, határidőkkel?
- Rögzítették a battle rhythm alapját és a verziókezelési szabályokat?

5.2.2. Phase 2 – Strategic Assessment (SSA): küldetés- és környezet-elemzés

A Phase 2 a közös helyzetértés és probléma-keretezés fázisa. A mission analysis során a törzs tisztázza a felhatalmazást (mandátum), a politikai és jogi korlátokat (pl. ROE, nemzeti caveat), a rendelkezésre álló erőforrásokat és a kívánt végállapotot (end state). A COPD szerint az end state nem pusztán „szándék”, hanem mérhető állapotok együttese, amelyhez success criteria tartozik. A success criteria feladata, hogy a művelet lezárásának és tranzíciójának alapját adja, és hogy később az assessmentben egyértelműen értékelhető legyen, teljesült-e a cél [33].

A Phase 2-ben a J2 hírszerzési támogatása kulcsfontosságú. A törzsnek meg kell értenie az ellenség szándékát, képességeit és sebezhetőségeit, valamint a legvalószínűbb és legveszélyesebb ellenséges COA-kat. A helyzetértés nemcsak katonai dimenzióban történik: a COPD az OE rendszerként való leírását ösztönzi, amelyben politikai, gazdasági, társadalmi és információs tényezők is megjelennek. A mission analysis kimenetének explicitnek kell lennie: a korlátozások és feltételezések dokumentált listája, kockázatregiszter vázlata, és a legfontosabb parancsnoki kérdések (commander’s critical questions) rögzítése. A legnagyobb buktató a túl korai megoldás: ha a törzs COA-t fejleszt, mielőtt a probléma modellje stabil, a későbbi tervezés rossz alapra épül [33], [34].

A Phase 2-ben célszerű az első operációs megközelítés (initial operational approach) vázlata: egy rövid narratíva arról, hogyan érhető el az end state. Ez még nem COA, hanem gondolkodási keret, amely segít kijelölni a design elemeket (lehetséges COG, előzetes DC-k, lehetséges LoE/LoO). Az initial approach és a feltételezéslista együtt mutatja meg, mekkora a bizonytalanság. Gyorsított tervezésben az initial approach „tartósan vázlat” maradhat, és a Phase 5 adaptációban finomodik; normál ütemben viszont a Phase 3-ban COA-kká alakul [33].

Ellenőrző kérdések (checklist):

- A problem statement, end state és success criteria egyértelmű és mérhető?
- A korlátozások és feltételezések listája elkészült és felelősökkel ellátott?
- Készült ellenséges COA elemzés (most likely/most dangerous)?
- A gap/CCIR lista frissült, és a gyűjtési prioritás ehhez igazodik?
- Van kezdeti kockázatregiszter és megfogalmazott risk appetite?

5.2.3. Phase 3 – MRO/COA-fejlesztés: alternatívák létrehozása és tesztelése

A Phase 3 célja a lehetséges katonai válaszopciók (MRO) és a cselekvési változatok (COA) kidolgozása. A COPD szerint az opciók akkor segítik a parancsnoki döntést, ha valóban eltérnek egymástól a megközelítésben és a kockázat–hatás profilban. Egy COA eltérhet például a main effort fókuszában, a tempo és erőkoncentráció mértékében, a partner erők bevonásának arányában, vagy a multidomain eszközök (cyber, information, fires) súlyozásában. Ha a COA-k csak „átfogalmazott” verziók, a döntési pont értelme megszűnik, mert a parancsnok valójában nem alternatívák között választ [33].

A COA-fejlesztés minimumleírása a COPD-ben: (1) operational approach röviden, (2) main effort és az erők fő irányai, (3) phasing és critical events, (4) kulcsfeladatok és előfeltételek (különösen LOG és CIS), (5) kockázatregiszter COA-specifikus elemei, (6) civil és információs következmények vázlata, valamint (7) kezdeti assessment logika. A Phase 3-ban a parancsnok számára különösen fontos, hogy az opciók „összemérhetők” legyenek; ezért a törzs jellemzően COA-sablonban dolgozik, és COA-összehasonlító mátrixot készít [33], [10].

A Phase 3 egyik kulcslépése a korai wargaming. A COPD a wargaminget nem a tökéletes előrejelzés eszközének tekinti, hanem a hibák és hiányok korai feltárása módszerének. Az action–reaction–counteraction (A-R-C) ciklus segít a COA logikai tesztelésében: mit teszünk, hogyan reagál az ellenség, milyen ellenreakcióval válaszolunk. A wargaming eredményei tipikusan: (a) gap-ek és képességhiányok; (b) döntési pontok és CCIR; (c) szinkronizációs problémák; (d) branch/sequel opciók körvonalai; valamint (e) kockázatok és mitigációs javaslatok. A wargaming dokumentálása (record sheet) azért kritikus, mert a későbbi COA-összehasonlítás csak így támasztható alá bizonyítékkal, nem csupán benyomásokkal [33].

Ellenőrző kérdések (checklist):

- Van legalább két-három érdemben eltérő COA, eltérő risk/impact profillal?
- Minden COA-hoz készült erő/idő/LOG/CIS előfeltétel-vizsgálat?
- Készült korai wargaming red teammel, dokumentált outputokkal?
- A kockázatregiszter és feltételezések COA-specifikusan frissültek?
- A COA-összehasonlítás bizonyítékokra támaszkodik, nem véleményekre?

5.2.4. Phase 4 – CONOPS és OPLAN/OPORD: integrált részletezés és végrehajthatóság

A Phase 4-ben a kiválasztott COA koncepcióból végrehajtható terv lesz. A COPD gyakorlati okból különválasztja a CONOPS (konceptió) és az OPLAN/OPORD (részletes terv) kidolgozását. A CONOPS a design formalizálása: rögzíti a COG/DC logikát, a LoE/LoO struktúrát, a phasinget, a main effort-ot, és a kezdeti assessment keretet. A CONOPS jóváhagyása minőségi kapu: ha a logika hibás vagy nem mérhető, a részletezés csak a hibát mélyíti. Az OPLAN/OPORD ezután a végrehajtás feltételeit biztosítja: feladat-hozzárendelés, szinkronizáció, erő- és képességrendelés, annexek integrációja, valamint a command and control és sustainment részletezése [33], [34].

A Phase 4-ben a decision support termékek kialakítása különösen fontos. A wargaming eredményeiből a törzs előre azonosítja a várható döntési helyzeteket (decision points), és ezekhez triggeret (indikátor + küszöb), CCIR-t és ajánlott opciókat rendel. Ez a Decision Support Template (DST) és a Decision Support Matrix (DSM) magja. A DST/DSM célja, hogy a Phase 5-ben a parancsnok ne ad hoc módon, hanem előre elemzett helyzetekben döntsön. Ez gyorsítja a döntési ciklust és csökkenti a kognitív terhelést, különösen akkor, amikor a művelet több domainben párhuzamosan zajlik. A COPD hangsúlyozza, hogy a DST/DSM csak akkor működik, ha az assessment indikátorok és küszöbök koherensen illeszkednek a DC-khez, és a CCIR valóban a döntéshez szükséges információra szűkül [33], [36].

A Phase 4 tipikus buktatója a késői szakági integráció. Ha a LOG, CIS, legal vagy information területek csak a végén „csatolnak” annexet, gyakori, hogy a terv belső ellentmondásos lesz: a műveleti tempo nem fenntartható, a C2 feltételek nem adóttak, vagy a jogi keret szűkebb, mint amit a targeting terv feltételez. A COPD ezért ösztönzi a korai integrációt: a szakági inputoknak már a COA- és CONOPS-szinten meg kell jelenniük, hogy a végrehajthatóságot a parancsnok idejében lássa. Ezzel csökkenthető a „késői visszabontás” jelensége, amikor a tervet a részletezés során kénytelenek újraépíteni [33], [34].

Ellenőrző kérdések (checklist):

- A CONOPS design koherens és jóváhagyott (COG/DC/LoE/LoO, phasing, main effort)?
- A DC-khez rendelték indikátorokat, küszöböket és adatforrásokat (assessment)?
- Az OPLAN/OPORD feladat-hozzárendelése összhangban van a LOG/CIS/legal annexekkel?
- A DST/DSM és CCIR döntési pontokhoz kötve elkészült és érthető?

- A terv tartalmaz branch/sequel opciókat és FRAGO traceability szabályt?

5.2.5. Phase 5 – Végrehajtás, assessment és FRAGO (adaptáció)

A Phase 5 a végrehajtás és folyamatos adaptáció fázisa. A COPD szemléletében a tervezés és végrehajtás nem különül el élesen: a terv baseline-t ad, de a környezet és az ellenség alkalmazkodása miatt a parancsnoknak rendszeresen módosítania kell a tevékenységet. A siker kulcsa, hogy a Phase 4-ben mennyire sikerült előre beágyazni a döntési pontokat és az assessment logikát. Ha a DST/DSM jól működik, a törzs a megfelelő információt a megfelelő időben tudja biztosítani, és a parancsnok előre azonosított opciók közül választhat. Ha nincs előretervezett döntéstámogatás, a döntés ad hoc lesz, a kockázatvállalás kevésbé kontrollált, és a FRAGO-k láncolata szétszedi a terv koherenciáját [33], [36].

A Phase 5 alatt különösen felértékelődik az információs dimenzió és a civil hatások kezelése. Egy reputációs incidens, egy félreértelmezett narratíva vagy egy kibertámadás percek-órák alatt stratégiai következményeket generálhat. A COPD ezért a fókuszát hangsúlyozza: CCIR, priorizált indikátorok és rövid, döntésorientált briefek. Az assessmentnek trendeket kell mutatnia, nem csak „állapotot”. A gyakorlatban ez azt jelenti, hogy a reporting cycle-ben a kulcsindikátorok (MOE) időbeli alakulása, a küszöbök átlépése és a várható következmények legyenek az elsődlegesek. A túl sok indikátor és a túl részletes jelentés paradox módon lassítja a döntést, mert a parancsnok nem kap világos képet arról, mi változott a döntési helyzetben [36].

A Phase 5-ben a feltételezéslista és a risk register frissítése döntő. A végrehajtás során derül ki, mely feltételezések voltak túl optimisták, és hol voltak a terv vakfoltjai. A COPD logikája szerint a kockázatmenedzsment nem egyszeri „beírás”, hanem folyamatos ciklus: azonosítás, értékelés, mitigáció és újraértékelés. A FRAGO-k esetén különösen fontos a traceability: a változtatás indoka, várható hatása és kockázata legyen dokumentált, különben a végrehajtás közben a döntések láncá válik átláthatatlanná. Ez nemcsak operatív, hanem jogi és politikai elszámoltathatósági kérdés is [33].

Ellenőrző kérdések (checklist):

- Az assessment ciklus döntést támogat és battle rhythm-be integrált?
- A risk register és feltételezéslista élő dokumentumként frissül?
- A FRAGO-k indoka, hatása, kockázata dokumentált (traceability)?
- A CCIR és trigger-küszöbök működnek (DST/DSM)?

- A lessons learned megfigyelések gyűjtése folyamatosan történik?

5.2.6. Phase 6 – Transition/Termination: kilépés, átadás és szervezeti tanulás

A Phase 6 a művelet lezárásának, átadásának és a hatás fenntartásának fázisa. A COPD szerint a kilépés nem utólagos adminisztráció, hanem a művelet tervezett része: a success criteria és az end state alapján kell meghatározni, mikor teljesültek az exit feltételek, és milyen ütemben adható át a felelősség. A tranzíció tipikusan több szereplő között történik (host nation, partner erők, civil szervezetek), ezért a felelősségek és függőségek rögzítése kritikus. Ha az exit criteria nem mérhető, a kilépés könnyen politikai „érzés” alapján történik, miközben a műveleti kockázatok (pl. instabilitás, erőszak visszatérése) nem kezelhetők. A mérhető tranzíciós kritériumok ezért ugyanúgy design és assessment kérdések, mint bármely más DC [33], [36].

A Phase 6 a lessons learned intézményesítésének elsődleges terepe. A JALLC kézikönyv szerint a szervezeti tanulás akkor történik meg, ha az observation–analysis–action plan–validation lánc végigfut, és a változtatás beépül a szervezet működésébe (SOP, sablon, képzés, eszköz). A COPD-termékek standardizáltsága lehetővé teszi, hogy a tanulságok konkrét dokumentumrészekhez kapcsolódjanak: például a risk register mezőinek módosítása, a wargame record sheet kötelező bevezetése, vagy az assessment indikátor-katalógus finomítása. A tanulságok intézményesítése csökkenti a következő művelet tervezési ciklusidejét és javítja a döntési minőséget, mert a korábbi hibákból származó „védőkorlátok” beépülnek a folyamatba [35].

Ellenőrző kérdések (checklist):

- A tranzíciós/exit kritériumok mérhetőek és illeszkednek a success criteria-hoz?
- Az átadás felelősei, üteme és fenntarthatósági feltételei rögzítettek?
- Készült final assessment és post-operation review?
- A lessons learned intézkedési terv validált (action plan + validation)?
- A tanulságok beépültek sablonokba, SOP-kba és képzésbe (institutionalization)?

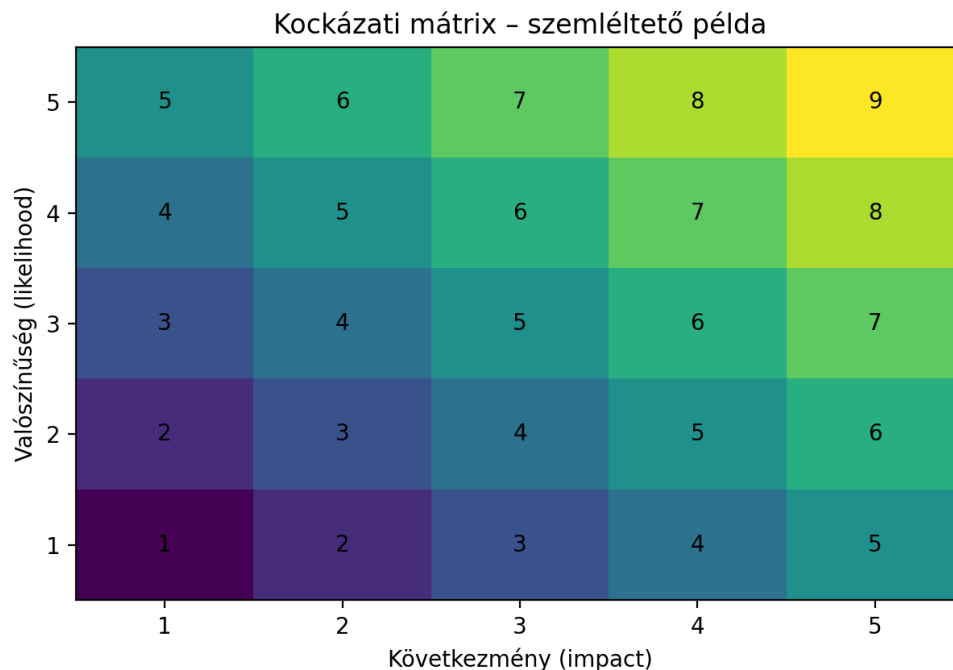
5.2.7. Gyorsított tervezés (accelerated planning) – a sebesség és bizonytalanság kezelése

A COPD elismeri, hogy válsághelyzetben gyakran nincs idő a teljes fázislogika „tankönyvi” végigfuttatására. A gyorsított tervezés célja nem a módszertan elhagyása, hanem

a párhuzamosítás és a fókusz. Tipikus megoldás, hogy bizonyos tevékenységek átfednek: a Phase 2 helyzetértéke közben már indul a Phase 3 COA-vázlatolás, miközben a J2 tovább pontosítja az ellenség modelljét. A gyorsítás ára az, hogy több feltételezést kell elfogadni, ezért a feltételezéslista és kockázatregiszter kezelése, valamint a Phase 5 alatti adaptációs képesség felértékelődik. A COPD szemléletében a gyorsítás akkor „biztonságos”, ha a döntési csomag minimum elemei nem sérülnek: a probléma és end state rögzített, legalább két opció összehasonlítható, és legalább egy dokumentált wargame fut le [33].

A gyorsított tervezésben a battle rhythm átalakul: rövidebb, gyakoribb parancsnoki update-ek és döntési briefek jellemzőek, miközben a háttérben több párhuzamos munkacsoport fut. A döntéstámogatás szerepe ezért nő: a DST/DSM korai vázlata már a COA-szinten kialakulhat, majd a végrehajtás során finomodik. A gyakorlatban a gyorsított tervezés egyik bevált eszköze a moduláris OPLAN: egy stabil baseline (küldetés, C2, LOG alapok, ROE keret), amelyhez a részletek FRAGO-kkal épülnek hozzá. Ezzel a törzs elkerüli, hogy a bizonytalan részletek miatt „megálljon” a tervezés, miközben a baseline biztosítja a koherens végrehajtási keretet [33], [10].

A gyorsítás legnagyobb veszélye a dokumentálatlanság. Időnyomás alatt könnyű úgy dönteni, hogy a feltételezéseket vagy az elvetett opciók indokait nem írják le. Ez azonban később súlyos problémát okoz: a Phase 5 adaptációban nem világos, mely döntések mely feltételezésre épültek, és a lessons learned is nehezebbé válik. A COPD logikájából levezethető jó gyakorlat, hogy gyorsított módban a dokumentumok terjedelmét csökkentjük, de a struktúrát megtartjuk: a döntési csomag legyen rövid, de tartalmazza a „kérdés–opció–kockázat–ajánlás” magot, és legyen változásnapló. Ez biztosítja az auditálhatóságot és a koherens végrehajtást [33]. Az 5.1. ábra egy tipikus kockázat elemzést szemléltet a bekövetkezési valószínűség és a hatás vonatkozásában, hozzárendelt pontértékkel.



Ábra 5-3. Kockázati mátrix szemléltető példa (saját szerkesztés) [33].

5.3. Műveletkialakítás (operational design): COG, LoO/LoE, hatások, kockázat

A COPD az operational design-t a tervezés módszertani magjaként kezeli. A design feladata, hogy a politikai célokat és korlátokat olyan műveleti logikává alakítsa, amely egyszerre koherens és adaptív. Koherens, mert a cél–hatás–feladat lánc belső ellentmondások nélkül összeáll. Adaptív, mert a terv előre kezeli a bizonytalanságot (feltételezések, kockázatregiszter, döntési pontok, branch/sequel). A design tehát nem díszítő ábrák összessége, hanem a végrehajtás során használt gondolkodási és döntési keret, amelyhez a törzs folyamatosan visszanyúl a FRAGO-k készítésekor [33], [34].

5.3.1. COG-elemzés és a CG–CC–CR–CV lánc gyakorlati értelme.

A COG (centre of gravity) a rendszer azon eleme, amelynek megőrzése vagy semlegesítése aránytalan hatással van a rendszer működésére. A Strange-féle megközelítés a COG-elemzést operacionalizálja: ha azonosítottuk a COG-t, meg kell határozni a kritikus képességeket (CC), amelyekre a COG támaszkodik; a kritikus követelményeket (CR), amelyek ezekhez a képességekhez szükségesek; és a kritikus sebezhetőségeket (CV), ahol beavatkozással aránytalan hatás érhető el. A tervezésben a CV-k gyakran „target system” jellegű beavatkozási

pontokká válnak, de csak akkor, ha a jogi/politikai korlátok és a másodrendű hatások figyelembe vannak véve [38].

5.3.2. *Decisive conditions (DC) és objectives (OOs).*

A DC-k a kívánt end state felé vezető úton elérendő állapotok. A DC nem feladat, hanem olyan állapot, amelyet a műveleti környezetben létre kell hozni vagy meg kell őrizni. A feladatok (tasks) DC-ket szolgálnak, a DC-k pedig objectives-hez és végül az end state-hez kapcsolódnak. A DC megfogalmazása akkor jó, ha mérhető és „operacionalizálható”: hozzá rendelhetők hatások (effects), feladatcsomagok és indikátorok. Tipikus hiba, amikor a DC túl általános, ezért nem mérhető, vagy túl feladat-szerű, ezért nem állapotként kezelik. A COPD a traceability elvvel segít: ha egy DC nem köthető indikátorhoz és feladatokhoz, akkor valószínűleg rosszul van megfogalmazva [33].

5.3.3. *LoO/LoE és szinkronizáció.*

A lines of operation (LoO) és lines of effort (LoE) a design vizualizációs és menedzsment eszköze. A LoE-k gyakran különböző funkcionális területeket és domain-eket kötnek össze: például security, governance, information, sustainment. A LoE akkor működik, ha tulajdonosa (LoE owner) van, és a LoE mentén rendszeres szinkron fórum működik, amely összehangolja a feladatokat és kezeli a függőségeket. A multidomain környezetben a LoE-k közötti kölcsönhatások különösen erősek: egy kinetikus művelet információs következményei visszahatnak a politikai mozgástérre, a cyber tevékenységek pedig a C2 és ISR működést befolyásolják. A design feladata, hogy ezeket a kapcsolatokat láthatóvá tegye, és a tervezés során megfelelő mitigációt (branch/sequel) építsen be [33], [36].

5.3.4. *Hatásláncok (effects tree) és mérhetőség.*

A hatások gyakran hierarchikusak: egy magas szintű DC elérése több köztes hatásból és feladatcsomagból áll. A hatáslánc segít abban, hogy az assessment ne csak „végállapot” szinten mérjen, hanem köztes jeleket (leading indicators) is figyeljen. Például egy információs LoE-ben a MOE lehet a bizalom trendje, de a leading indicator lehet a narratív térben mért ellenséges aktivitás vagy a médiaelérés. A leading és lagging indikátorok kombinációja csökkenti a későn reagálás kockázatát. A commander's assessment kézikönyv hangsúlyozza, hogy az indikátorok

számát korlátozni kell: a döntéstámogatás szempontjából a legfontosabb az, ami döntést vált ki (küszöbök), nem az, ami érdekes adat [36].

5.3.5. Kockázatkezelés: risk register, risk appetite és mitigáció

A COPD a kockázatkezelést a döntési csomag részének tekinti. A risk register a kockázat leírásán túl rögzíti: a kiváltó okot (gyakran feltételezés), a várható következményt, a valószínűség/hatás értéket, a mitigációs lépéseket, a triggereket és a felelőst. A döntési pontokon a parancsnok valójában risk appetite-t rögzít: mekkora kockázat vállalható a cél érdekében. A mitigáció gyakran branch/sequel formájában jelenik meg: ha egy trigger teljesül, előre megtervezett váltást hajtunk végre. A risk register így összeköti a design-t (mit akarunk elérni) a decision support-tal (mikor kell váltani) és az assessmenttel (miből látjuk, hogy váltani kell) [33], [36].

Az 5-2. táblázat egy részletes risk register sablont és szemléltető bejegyzéseket tartalmaz.

ID	Kockázat	Ok/feltételezés	Hatás	L/I	Mitigáció	Trigger/DP	Owner
R-01	Hiányos közös helyzetkép	Besorolási korlátok	Torz SSA/COA	4/3	releasable összefoglaló; alternatív indikátorok	DP-1	J2/J5
R-02	LOG culminatio n túl korán	optimista utánpótlás	tempo csökken	3/5	tartalék; útvonal redundancia; HNS	DP-4	J4
R-03	CIS kiesés/zavarás	redundancia hiánya	C2 sérül	3/5	hardening; alternatív hálózat	DP-3/4	J6
R-04	IO reputációs kár	civil hatás + ellenség IO	politikai támogatás csökken	3/5	STRATCOM terv; monitoring; ROE review	DP-A	STRATCOM/ Legal
R-05	Ellenség gyors adaptáció	wargame alulbecsli	DC nem	4/4	red teaming;	DP-2/5	J2/J3

			teljesül		branch/sequel		
R-06	Partner kapacitás gyenge	képességhiány	transit ion kockázat	4/4	capacity building; phased handover	DP-6	J5/J9

5-2. táblázat. Risk register (minta; saját szerkesztés) [33], [34].

5.4. COA-fejlesztés, wargaming és összehasonlító értékelés

A COPD a COA-fejlesztést a döntési minőség kulcsának tekinti, mert itt válik kézzelfoghatóvá a parancsnok számára a kockázat–hatás kompromisszum. A COA leírásának minimuma: operational approach, main effort, phasing, critical events, erő- és fenntarthatósági előfeltételek (LOG/CIS), civil és információs következmények, valamint kockázat. A COA-k összehasonlíthatósága érdekében a törzs tipikusan COA-sablonban dolgozik, és többkritériumos összehasonlító mátrixot használ. A JP 5-0 hasonlóan rögzíti, hogy a COA-k értékelése a suitability, feasibility és acceptability dimenziók köré épül [10].

5.4.1. Wargaming és red teaming.

A wargaming célja a COA tesztelése az ellenség és a környezet reakcióinak figyelembevételével. Az action–reaction–counteraction ciklus segít a logikai hiányok feltárásában. A red team szerepe különösen fontos, mert időnyomás alatt a törzs hajlamos a megerősítési torzításra: azt keresi, ami a preferált COA-t igazolja. A red team feladata, hogy az ellenség logikáját következetesen képviselje, és a COA gyenge pontjait „kényszerítse” felszínre. A wargaming outputjai: döntési pontok, CCIR, gap-ek, kockázatok, mitigáció, valamint branch/sequel opciók. Ezek képezik a döntéstámogató termékek (DST/DSM) alapját [33].

5.4.2. COA-összehasonlítás és döntési transzparencia.

A pontozás csak akkor értelmes, ha a pontszám mögött bizonyíték van: wargame esemény, LOG számítás, jogi elemzés vagy hírszerzési értékelés. A COPD elvárja, hogy a választott COA indoka és az elvetett opciók elutasításának oka is dokumentált legyen. Ez a Phase 5-ben

is hasznos: ha a végrehajtás során a feltételek változnak, és felmerül a COA-váltás, a törzs gyorsan vissza tud nyúlni az eredeti összehasonlításhoz, és nem kell a vitát „nulláról” újraindítani [33].

5.4.3. Decision Support Template/Matrix.

A DST és DSM a wargaming eredményeit döntési formába önti: döntési kérdés, trigger, CCIR, opciók és ajánlás. A DSM idő- és térbeli kontextusba helyezi a döntési pontokat. A döntéstámogatás akkor működik, ha a triggerek a design elemekhez (DC-khez) kapcsolódnak, és az assessment indikátorok valóban mérhetőek. A túl sok döntési pont és túl sok CCIR azonban megterhelheti a reporting rendszert; ezért a COPD fókuszra kér: a DST/DSM-ben csak a valóban parancsnoki döntést igénylő helyzetek szerepeljenek [33], [36].

Kritérium	Súly	COA-A	COA-B	COA-C	Bizonyíték
Suitability					
Feasibility (LOG/CIS)					
Acceptability (risk)					
Political/Legal					
Civil/Info effects					
Sustainability					
Interoperability/C2					
Flexibility (branch/sequel)					
Deployability/tempo					

5-3. táblázat. COA összehasonlító mátrix sablon; (saját szerkesztés) [33], [10].

Idő/fázis	Saját akció	Ellenség reakció	Ellenreakció	DP/CCIR	Kockázat/gap

5-4. táblázat. Wargame record sheet (sablon; saját szerkesztés) [33].

DP	Kérdés	Trigger	CCIR	Opciók	Ajánlás
----	--------	---------	------	--------	---------

DP-A	Reputációs incidens kezelése szükséges?	Incidensek>X/48h	civil casualties; OSINT trend	FP erősítés; IO kampány; ROE review	IO+FP
DP-B	Ellenség C2 adaptáció miatt váltani kell?	SIGINT minta változik	redundancia jelei	targeting váltás; cyber priorítás	cyber+targeting
DP-C	LOG késések miatt tempo módosul?	készletszint<Y%	útvonal biztonság	tempo csökkentés; útvonalváltás	tempo+HNS

5-5. táblázat. *Decision Support Template (minta; saját szerkesztés) [33], [36].*

5.5. Műveletértékelés (assessment), mérőszámok, lessons learned beépítése

A COPD az assessmentet a parancsnoki döntési ciklus részeként kezeli. Az assessment célja nem a jelentéskészítés, hanem a döntéstámogatás: trendeket kell azonosítani, értelmezni a hatásokat, és javaslatot adni arra, hogy a terv maradjon-e érvényben vagy módosítás szükséges. Ezért a tervezés során már a CONOPS/OPLAN kidolgozásakor rögzíteni kell az assessment logikai modelljét (tasks → outputs → effects/DC → objectives → end state), a kulcsindikátorokat és a küszöböket. A Commander's Handbook for Assessment Planning and Execution hangsúlyozza, hogy az indikátorok számát korlátozni kell, és a reporting cadence-et a döntési pontokhoz kell igazítani [36].

5.5.1. MOP és MOE, leading és lagging indikátorok.

A measure of performance (MOP) azt méri, hogy a feladatot végrehajtottuk-e, a measure of effectiveness (MOE) pedig azt, hogy közelebb jutottunk-e a célhoz. A MOE gyakran indirekt, ezért proxy indikátorokra és triangulációra van szükség. A leading indikátorok korai jelek, amelyek előre jelzik a trendet (pl. ellenség aktivitási minta), míg a lagging indikátorok a tényleges eredményt mutatják (pl. bizalom, incidensszám). A két típus kombinációja csökkenti a későn reagálás kockázatát, és segít a DST/DSM triggerek működtetésében [36].

5.5.2. Data governance és besorolási korlátok.

Az assessment több adatforrásból dolgozik (ISR, jelentések, OSINT, civil szervezetek). A kihívás az adatminőség és a besorolás kezelése. A COPD logikájából levezethető jó gyakorlat, hogy a koalíciós döntésekhez kialakítunk egy közös, releasable indikátorkészletet, és ettől elkülönítjük a szenzitív indikátorokat. A data governance része a validáció és a torzításkezelés: mely forrás mennyire megbízható, milyen szabály szerint kombináljuk a forrásokat, és hogyan dokumentáljuk az értelmezési bizonytalanságot (confidence level) [33], [36].

5.5.3. Lessons learned intézményesítése.

A JALLC kézikönyv szerint a szervezeti tanulás akkor sikeres, ha a tanulságokból intézkedési terv lesz, és a változtatást validálják. A COPD standardizált termékei lehetővé teszik, hogy a tanulságokat konkrét sablonokba és SOP-kba építsük be: például a risk register mezőinek finomítása, a wargame record sheet kötelező alkalmazása, vagy az assessment indikátor-katalógus kialakítása. A tanulságok beépítése csökkenti a következő tervezési ciklus idejét és növeli a döntési minőséget, mert a korábbi hibákból származó „védőkorlátok” beépülnek a folyamatba [35].

Elem	Tartalom	Kapcsolat
Logikai modell	tasks→outputs→effects/DC→OOs→end state	CONOPS/OPLAN
Indikátor-katalógus	MOP/MOE + küszöbök + trend	DST/DSM triggerek
Adatgyűjtési terv	forrás, felelős, gyakoriság	J2/J3 cycle
Adatminőség	validáció, trianguláció, bias	confidence level
Reporting cadence	brief ritmus és forma	battle rhythm
Döntési integráció	mely DP-hez milyen assessment kérdés	FRAGO/branch

5-6. táblázat. Assessment terv minimális felépítése (saját szerkesztés) [36].

5.6. Jelen kihívások a NATO hadműveleti tervezésben (időnyomás, adatbőség, komplexitás)

A COPD v3.1 alkalmazása olyan környezetben történik, ahol egyszerre nőtt a gyors reagálás igénye és a műveleti környezet komplexitása. Az időnyomás gyakran kényszeríti a gyorsított tervezést; az adatbőség és dezinformáció csökkenti a jel–zaj arányt; a multidomain műveletek pedig a szinkronizációt és interoperabilitást nehezítik. Ezek a tényezők közvetlenül növelik a kognitív terhelést, és erősítik a torzítások kockázatát. A COPD kezelési logikája éppen ezért a fókuszra és döntési transzparenciára épül: CCIR, döntési pontok, kockázatregiszter és traceability. Minél gyorsabb és komplexebb a környezet, annál nagyobb értéke van a strukturált tervezési fegyelemnek [33].

5.6.1. Időnyomás és *accelerated planning*.

Gyorsított tervezésben a legfontosabb kockázat, hogy a törzs kihagy minőségi kapukat. A COPD szerint akkor is szükséges rövid mission analysis brief, legalább két opció összevetése, és dokumentált wargaming. A gyorsítás ára a bizonytalanság, ezért a feltételezéslista, a risk register és a DST/DSM felértékelődik. A gyakorlatban a gyorsított tervezés sikerének feltétele a döntési minimum megőrzése: rövid dokumentumok, de megmaradó struktúra és változásnapló [33], [10].

5.6.2. Adatbőség, OSINT és dezinformáció.

A modern műveletekben a szenzorok és nyílt források hatalmas adatmennyiséget generálnak, miközben az ellenség aktívan manipulálja az információs teret. A döntéshez releváns információ kiválasztása ezért kulcskérdés. A COPD eszköze a CCIR és a priorizált indikátor-készlet: a törzs nem mindent gyűjt, hanem azt, ami a döntési pontokhoz kell. Az assessmentben a trianguláció és a confidence level rögzítése csökkenti a manipuláció kockázatát [36].

5.6.3. Multidomain komplexitás és interoperabilitás.

A több domainben zajló műveletekben a hatások egymásra rétegződnek és nemlineárisak. A COPD a design (DC/LoE), a wargaming és a branch/sequel eszközeivel segít kezelni az

adaptációt. Interoperabilitási korlátok esetén a terminológiai fegyelem és a termékek releasability-rétegezése csökkenti a félreértéseket, és gyorsítja a koalíciós integrációt. A kihívások kezelésében a lessons learned intézményesítése is fontos, mert a gyors környezetben a tanulságok beépítése versenyelőnyt jelent [35], [37].

5.7. Összefoglalás

A fejezet rámutatott, hogy a COPD v3.1 a NATO hadműveleti tervezés végrehajtás-közelí eljárásrendje, amely a doktrinális AJP-5 elveket a törzsek gyakorlati munkájára fordítja le. A COPD erőssége a fázislogika, a döntési pontok, a termékrendszer és a traceability elv együttese, amelyek biztosítják, hogy a tervezés döntéstámogató, auditálható és adaptív legyen. A tervezés minőségét különösen a problémakeretezés (SSA), az alternatívák valódi diverzifikációja (MRO/COA), a wargaming dokumentáltsága, valamint az assessment és decision support integrációja határozza meg [33], [34], [36].

A COPD strukturált folyamata és termékei olyan „rögzített pontokat” adnak, amelyekhez technológiai (pl. MI-alapú) támogatás illeszthető. A gap/CCIR karbantartása, a össz-adatforrású helyzetkép szintézise, a COA-variánsok gyors előállítás és a wargame outputok strukturált rögzítése mind olyan terület, ahol gépi eszközök csökkenthetik a ciklusidőt és a kognitív terhelést. Ugyanakkor a COPD logikája egyértelműen kijelöli a határt: a kockázatvállalás és a döntési felelősség parancsnoki kérdés, ezért minden automatizált támogatásnak transzparensnek, auditálhatónak és ember által felügyeltnek kell lennie.

A tervezési sablonokat és kitölthető űrlapokat (COPD-kompatibilis), valamint a tervezési segédleteket az A-I mellékletek tartalmazzák.

6. FEJEZET MI TÁMOGATÁS A COPD 1-2 FÁZISOKBAN: A MŰVELETI KÖRNYEZET MEGÉRTÉSE, STRATÉGIAI HELYZETÉRTÉKELÉS (INITIAL SITUATIONAL AWARENESS, STRATEGIC ASSESSMENT)

A NATO Comprehensive Operations Planning Directive (COPD) a stratégiai és műveleti szintű tervezés egyik központi hivatkozási rendszere. A COPD szerint az 1. fázis (Initial Situational Awareness) azonosítja és értelmezi a lehetséges vagy tényleges válságot, és előkészíti a további tervezést; a 2. fázis (Strategic Assessment) pedig SACEUR Stratégiai Helyzetértékelésének (SSA) kidolgozását és koordinációját szolgálja [33].

A két fázis közös kihívása, hogy a döntési ablak rövid, miközben a műveleti környezet (Operating Environment, OE) sokdimenziós és gyorsan változó. Az AJP-5 hangsúlyozza, hogy az OE megértése kritikus előfeltétele a tervezésnek, és a JIPOE a PMESII spektrum lefedésével támogatja a holisztikus helyzetértést [34].

A mesterséges intelligencia (MI) ebben a korai szakaszban nem önálló döntéshozó, hanem képességnövelő: gyorsítja a össz-adatforrású információk feldolgozását, támogatja a strukturált elemzést (indikátorok, trendek, ok-okozati összefüggések), és javítja a helyzetkép frissíthetőségét.

A fejezet célja, hogy a COPD 1-2 fázisaihoz illeszkedő, MI-vel támogatott elemzési és adatkezelési megközelítést mutasson be a DIME/PMESII és kapcsolódó keretek, a JIPOE, a össz-adatforrású felderítés, az adatfűzió, a dezinformáció/deception kezelése, a prediktív fenyegetésvértékelés és az alkalmazható MI-eszközkészlet bemutatásával.

6.1. A DIME/PMESII és kapcsolódó elemzési keretek (ASCOPE/ICR2) – adatmodell és indikátorok, a PMESII, JIPOE és össz-adatforrású felderítés adatforrásai

A COPD 1-2 fázisaiban a törzs feladata a műveleti környezet strukturált leírása és az ebből levezethető stratégiai következtetések megfogalmazása. Ez két nézőpontot kapcsol össze: a környezet rendszerszintű megértését (PMESII), valamint a lehetséges beavatkozási eszközöket (DIME).

A PMESII (political, military, economic, social, infrastructure, information) a megfigyelés és helyzetértés tengelye. A DIME (diplomatic, informational, military, economic) a lehetséges befolyásolási és cselekvési opciók tengelye. A kettő együtt segíti, hogy a helyzetértékelés közvetlenül kapcsolódjon a válaszopciókhoz.

A PMESII-PT változat a fizikai környezet és az idő dimenziójával egészíti ki a keretet, így a térbeli korlátok, az időkritikusság és a szezonális tényezők expliciten kezelhetők [35].

Az ASCOPE bontás (areas, structures, capabilities, organizations, people, events) a PMESII dimenziókat a tér, az infrastruktúra, a szervezetek és a populáció szintjére „lefordítja”. Az ICR2 jellegű keretek pedig a kapcsolatok, kommunikációs csatornák és szabályok mentén segítik a társadalmi és információs dinamika megértését, ami különösen fontos a hibrid fenyegetések esetén.

A PMESII-PT és az ASCOPE kombinálása a gyakorlatban gyakran egy „mátrix” formájában jelenik meg: a PMESII dimenziók (P–M–E–S–I–I) a rendszer-szintű környezeti tényezőket, míg az ASCOPE kategóriák (A–S–C–O–P–E) a konkrét terepi megjelenéseket (területek, struktúrák, képességek, szervezetek, emberek, események) rendezik egymás mellé. A 6-1. táblázat a két keret metszetében szereplő tipikus példatényezőket mutatja be, amelyeket a stratégiai helyzetértékelésben indikátorokká és felderítési követelményekké lehet alakítani [35].

ASCOPE \ PMESII	Politikai (P)	Katonai (M)	Gazdasági (E)	Társadalmi (S)	Infrastruktúra (I)	Információ (I)
Területek (A)	Körzethatárak Párthovatartozási/ támogatott sági területek	Érdekeltségi terület (AOI) Befolyási és műveleti területek Biztonságos menedékek	Ágazatok Formális/informális gazdaság Természeti erőforrások Kereskedelmi útvonalak Áruk és szolgáltatások mozgása	Diaszpórák Táborok Enklávák Migrációs mintázatok Lakónegyedek Befolyási övezetek határai	Kereskedelmi Ipari Lakó Rurális Urbánus	Különböző médiatípusok lefedettsége

Struktúra (S)	Kormányzati központok	Bázisok Parancsnokságok (HQ)	Bankrendszer/pénzügy Üzemanyag Gyártás Raktározás Piacok Állatvásárok	Büntetés - végrehajtási intézetek Történelmi építmények és létesítmények Könyvtárak Iskolák és egyetemek Stadionok	Vészhelyzeti menedékek Energielosztás Földgáz Erőművek Egészségügyi létesítmények (kórházak) Középületek Közlekedés (repülőterek, hidak, kikötők, vasutak, utak és autópályák) Hulladékkezelés és -tárolás Vízkezelés és -tárolás	Kommunikációs vonalak Tornyok Vezetési és irányítási (C2) rendszerek Internetszolgáltatások Celluláris és mobil szolgáltatások Postai szolgáltatás Nyomdák Telefon
Képességek (C)	Adminisztráció Intézményi keretek (törvényhozás, igazságszolgáltatás stb.)	Doktrína Kiképzési anyagok Személyi állomány Létesítmények Felszerelés stb.	Fiskális képességek (valuta, monetáris politika)	Egészségügy Nyelv és dialektusok Társadalmi hálózatok Tulajdoni nyilvántartások és kontroll	Tiszta ivóvíz Ruházat Kommunikációs rendszer Rendfenntartás Tűzoltás Egészségügyi kapacitás Közegészségügy/szanitáció	Helyi kommunikációs hálózatok Más területekhez vezető kommunikációs kapcsolatok Internet-hozzáférés Nyomtatott anyagok Propagandamechanizmusok Rádió Televízió Üzenetküldés stb.
Szervezetek (O)	Nemzetközi partnerek Politikai pártok stb.	Határőrség Büntetés-végrehajtás Felkelők Terroristák Törzsi	Üzleti/vállalkozói szervezetek Céhek Szakszervezetek Önkéntes csoportok	Klán- és közösségi szerveződések Bűnözői és családi hálózatok	Kormányzati minisztériumok Üzleti közösségek	Médiacsoportok Befolyásos csoportok Külföldi kormányok Nemzetközi szervezetek csoportjai stb.

		milíciák stb.		Hazafias és szolgálat i szerveze tek Vallási szerveze tek Törzsek stb.		
Ember ek (P)	Vezetők Partnerség ek	Kulcsvez etők	Foglalkozási csoportok Üzleti vezetők Fogyasztási mintázatok Átlagos napi bér Jövedelemel oszlás	Oktatás Etnicitás Kulcssze replők Faji/etni kai szerkeze t Sérüléke ny csoporto k Nemi szerepek stb.	Tisztviselők Helyi lakosság Üzleti közösségek Létesítményvez etők Dolgozók stb.	Döntéshozók Vezetők Újságkiadók Újságírók Média
Esemé nyek (E)	Választáso k Találkozók Beszédék	Történel mi eseménye k Nem- harc eseménye k	Mezőgazdasá gi és piaci ütemtervek Betakarítási időszakok Piaci napok Fizetésnapok Vetési/ültetés i ütemterv	Ünnepsé gek Népszá mlálás Polgári zavargás ok Bűnügyi esemény ek Nemzeti ünnepek Vallási ünnepek	Karbantartási tevékenységek Projektek kivitelezése	Hírciklusok Sajtótájékoztató k Csoporttalálkoz ók

6-1. táblázat: PMESII–ASCOPE mátrix (példatényezők) (saját szerkesztés [35] alapján).

A PMESII-PT a katonai tervezésben széles körben alkalmazott kiinduló elemző eszköz: célja, hogy a műveleti környezetet (OE) több, egymással kölcsönhatásban álló „működési változó” mentén bontsa fel. A PMESII a politikai, katonai, gazdasági, társadalmi, információs és infrastruktúra dimenziókra utal; a később hozzáadott fizikai környezet és idő (PT) pedig azt biztosítja, hogy a terep–idő tényezők ne „mellékesen”, hanem a helyzetértékelés részeként jelenjenek meg [35].

A forrásanyag hangsúlyozza, hogy a módszertan amerikai hadsereg-beli (Army) törekvésekből nőtt ki: a haderő olyan heterogén műveleti helyzetekben tevékenykedik – a magas intenzitású konfliktustól a poszt-konfliktus stabilizáción át a humanitárius és katasztrófavédelmi műveletekig –, amelyekben a katonai komponens mellett civil intézményekkel, nemzetközi szervezetekkel és nem kormányzati szereplőkkel való koordináció is meghatározó [35]. Az OE feltételei gyorsan változhatnak, és gyakran összetettséggel, volatilitással, bizonytalansággal, instabilitással és kétértelműséggel írhatók le; a PMESII-PT nyolc változója ezért a „teljes kép” felépítését támogatja [35].

A PMESII-PT kialakításának kettős indoka különösen releváns a COPD 1–2 fázisában. Egyrészt a hagyományos konfliktusokban a célrendszer gyakran jól azonosítható (pl. kinetikus képességek), míg stabilizációs/államépítési környezetben a siker kulcsa sokszor a kormányzás, az ellátórendszerek, a gazdasági működés és az információs tér összefüggő rendszereinek megértése. Másrészt a közös tervezési közösségben az ASCOPE-t sok esetben minimumként alkalmazzák az információk begyűjtésére és strukturálására, de a PMESII-PT szükséges a rendszerszintű összefüggések finomításához és priorizálásához [35].

A két keret összekapcsolásának operatív eszköze a PMESII–ASCOPE mátrix. A mátrix logikája, hogy minden PMESII dimenzióban megnevezzük azokat a területi jellemzőket, struktúrákat, képességeket, szervezeteket, kulcsszereplőket és eseményeket, amelyek a műveletet érdemben befolyásolhatják. A metszet-cellák így nem „adatlisták”, hanem olyan paraméter-katalógusok, amelyekből (1) indikátorok, (2) hipotézisek, (3) hiánylisták és (4) felderítési kérdések vezethetők le. Ez közvetlenül támogatja a COPD 1–2 fázis stratégiai helyzetértékelését és a JIPOE-hoz illeszkedő információigény-képzést [35].

A táblázatos megjelenítés MI-támogatás szempontból is hasznos: a mátrix egy jól definiált taxonómiát ad a többforrású adatok címkézéséhez és összerendeléséhez. Például az NLP-alapú entitásfelismerés és szereplő–kapcsolat kinyerés a „Szervezetek/Emberek” cellákhoz köthető, míg a GEOINT/IMINT alapú objektum-detektálás és infrastruktúra-osztályozás a

„Struktúrák/Infrastruktúra” metszetben strukturálható. A mátrix így egyszerre támogatja a kvalitatív szakértői gondolkodást és az automatizálható adatfúziós munkafolyamatokat [35].

6.1.1. Adatmodell: entitások, kapcsolatok, események és indikátorok

A COPD 1-2 fázisban a feladat nem az, hogy minden lehetséges adatot összegyűjtsünk, hanem az, hogy a döntéshez releváns információt olyan szerkezetbe rendezzük, amelyben a bizonytalanság, az ellentmondás és a változás is kezelhető. Ehhez célszerű egy közös, többforrású adatmodellt alkalmazni.

A javasolt adatmodell négy alapkategoriat különít el: (1) entitások (szereplők, objektumok), (2) kapcsolatok, (3) események, (4) indikátorok. A komponensek együtt teszik lehetővé, hogy a PMESII/ASCOPE logika mentén gyűjtött adatokból ok-okozati következtetések és előrejelzések készüljenek.

- Entitások: állami és nem állami szereplők, szervezetek, kulcsszemélyek, fegyveres és civil csoportok, kritikus infrastruktúra elemek, képességek, információs csatornák.
- Kapcsolatok: parancsnoki és szervezeti viszonyok, finanszírozás, logisztikai láncok, kommunikációs kapcsolatok, szövetségek, rivalizálás, befolyásolási viszonyok, függőségek.
- Események: incidensek, katonai mozgások, politikai döntések, tüntetések, kibertámadások, információs műveletek, gazdasági intézkedések. Az események időbélyeggel és georeferenciával kerülnek rögzítésre.
- Indikátorok: megfigyelhető mintázatok vagy mérőszámok, amelyek változást jeleznek egy PMESII-dimenzióban, és amelyekhez elemzői hipotézis, küszöbérték és ellenőrzési terv rendelhető.

A COPD stratégiai helyzetértékelési folyamata kiemeli, hogy a törzsnek a stratégiai környezet (PMESII, METOC és geospatial) elemzése során azonosítania kell a kritikus információ- és tudáshiányokat, amelyek gyűjtési vagy beszerzési igényekké alakulnak [33]. Az indikátor-alapú megközelítés e hiányok célzott kezelését támogatja: meghatározza, mit kell figyelni, milyen forrásból, milyen gyakorisággal, és milyen változás esetén szükséges a helyzetértékelés frissítése.

6.1.2. DIME–PMESII összerendelés és indikátor-példatár

A DIME eszközszer és a PMESII dimenziók összerendelése segíti, hogy a helyzetértékelés közvetlenül kapcsolódjon a lehetséges válaszopciókhoz. A 6-2. táblázat illusztratív példákat ad a dimenziókhoz rendelhető indikátorokra és lehetséges DIME-hatásokra.

PMESII dimenzió	Indikátorok (példák)	Kapcsolódó DIME-reakciók (példák)
P – Politikai	kormányzati stabilitás; elit-frakciók törésvonala; jogszabályi/hatósági lépések; legitimáció és narratíva	D: diplomáciai koordináció; I: stratégiai kommunikáció; M: elrettentő jelenlét; E: célzott ösztönzők/szankciók
M – Katonai	erőkészültség; csapatmozgások; logisztikai készletek; integrált légvédelem aktivitása; harcrend-változás	M: erőátcsoportosítás; D: katonai együttműködés; I: elrettentő üzenetek; E: exportkontroll/finanszírozás
E – Gazdasági	árfolyam/infláció; energiaellátás; kritikus ellátási láncok; pénzügyi tranzakciós mintázatok; feketepiac	E: szankció/segély; D: gazdasági diplomácia; I: költség-narratíva; M: kulcsterületek védelme
S – Társadalmi	tüntetés-intenzitás; migráció; közbizalom; etnikai/vallási feszültségek; humanitárius igények	I: reziliencia; D: civil partnerség; M: védelem/támogatás; E: segélyprogram
I – Információs	domináns narratívák; bot-aktivitás; média-ökoszisztéma; platformszabályok; dezinformációs kampány jelei	I: StratCom és counter-disinfo; D: kommunikációs koordináció; M: IO integráció; E: platform-együttműködés
II – Infrastruktúra	energia, víz, egészségügy; közlekedési csomópontok; távközlés; kritikus létesítmények sérülékenysége	M: infrastruktúravédelem; E: helyreállítás; D: koordináció; I: bizalomépítés

6-2. táblázat: DIME–PMESII összerendelés és indikátor-példák (saját szerkesztés) [33], [34].

6.1.3. Adatforrások: PMESII, JIPOE és össz-adatforrású felderítés

A közös adatmodell csak akkor működik, ha a PMESII dimenziókhoz és az ASCOPE/ICR2 bontásokhoz hozzárendeljük a tipikus adatforrásokat és azok megbízhatósági jellemzőit. A JIPOE célja, hogy a tervező törzs számára holisztikus képet adjon a műveleti környezetről, és értékelje a környezet hatását a küldetés teljesítésére [34]. Ennek adatbázisát a multi-source felderítés és a kapcsolódó polgári adatok együttese adja.

A 6-3. táblázat egy gyakorlati forrásmátrixot mutat be: milyen PMESII területekhez melyik gyűjtési diszciplína és milyen adatcsomagok illeszthetők. A táblázat célja nem az exhaustivitás, hanem a tervező munka támogatása azzal, hogy a gyűjtési igények (RFIs, collection requirements) gyorsan lefordíthatók konkrét adatbeszerzésre.

PMESII	Primer adatforrások	Kiegészítő források	Jellemző MI-támogatás
Politikai	HUMINT, OSINT (kormányzati kommunikáció, elit-hálózat)	diplomáciai jelentések, szakértői adatbázisok	NLP: entitás- és kapcsolatkinyerés; narratíva-elemzés
Katonai	SIGINT, GEOINT/IMINT, MASINT	HUMINT, nyílt műholdképek	CV: objektum- és változástetektálás; anomáliaészlelés
Gazdasági	OSINT (makro és piac), CYBINT (pénzügyi csalás jelek)	kereskedelmi adatok, szankciólisták	idősor és hálózati elemzés; anomália a tranzakciókban
Társadalmi	OSINT (közösségi), HUMINT (lokális hangulat)	NGO és humanitárius adatok	NLP: témamodellezés; területi hotspot elemzés
Információs	OSINT, SIGINT (kommunikációs mintázatok), CYBINT	platform jelentések, fact-check hálózatok	bot-detektálás; tartalom-hitelesség; gráfelemzés
Infrastruktúra	GEOINT/IMINT, OSINT (üzemeltetői adatok)	műszaki szenzorok, jelentések	CV: sérülés/üzemállapot; digitális iker jellegű modellek

6-3. táblázat: Forrásmátrix PMESII dimenziók szerint (saját szerkesztés, JIPOE és multi-source logikával) [34], [36].

6.2. JIPOE folyamata és termékei; kapcsolódási pontok a COPD-hoz

A JIPOE (Joint Intelligence Preparation of the Operational Environment) olyan közös eljárásrend, amely a műveleti környezet strukturált elemzésével támogatja a parancsnok és a törzs döntéshozatalát. A JP 2-01.3 a JIPOE-t négy lépésben írja le: a műveleti környezet definiálása, a környezet hatásának leírása, az ellenség és más releváns aktorok értékelése, valamint az ellenség (és egyéb aktorok) lehetséges cselekvési változatainak (COA) meghatározása [10].

Az AJP-5 kiemeli, hogy a tervezéshez a parancsnoknak és a törzsnek közös, holisztikus képet kell kialakítania az OE-ről, amelyben a válság háttere, okai, a szereplők céljai és a dinamikák egységes narratívává állnak össze [34]. A JIPOE-hoz illesztett MI-képességek ezt a közös képet gyorsabb adatfeldolgozással, konzisztens tér-idő rögzítéssel és következetes indikátor-követéssel támogatják.

A JP 2-01.3 külön hangsúlyozza, hogy a JIPOE termékeit időben kell disszeminálni és integrálni a tervezésbe; ha a helyzet nem teszi lehetővé teljes írásos hírszerzési becslés elkészítését, akkor is a JIPOE sablonokat, mátrixokat, grafikonokat és más adatforrásokat el kell juttatni a törzs többi eleméhez, hogy azok beépüljenek a párhuzamos tervezési tevékenységbe [10]. Ez a gondolat közvetlenül kapcsolódik a COPD 1-2 fázis időkritikus jellegéhez.

A JIPOE lépései és termékei a COPD 1-2 fázisaiban két módon kapcsolódnak: (1) tartalmilag, mert a COPD stratégiai környezet elemzése PMESII, METOC és geospatial fókuszú, és kifejezetten információs hiányokat azonosít [33]; (2) folyamat-szinten, mert a JIPOE ciklikus frissítése és minőségbiztosítása a tervezés során folyamatos visszacsatolást biztosít.

6.2.1. A JIPOE lépések és a COPD 1-2 fázisok összekapcsolása

A 6-4. táblázat összefoglalja, hogy a JIPOE négy lépése milyen COPD-feladatokhoz és tipikus termékekhez rendelhető a korai tervezési szakaszban.

JIPOE lépés [10]	COPD 1. fázis (Initial Situational Awareness) [33]	COPD 2. fázis (Strategic Assessment) [33]	Tipikus termékek
1. OE definiálása	válságkeret kijelölése, AO/area scope, kulcsfogalmak és kezdeti feltételezések	AOI és releváns rendszerek pontosítása, hatókör-finomítás	OE leíró összefoglaló, alaptérképek, szereplőlista, adat-rendszerezési séma
2. OE hatásának leírása	gyors PMESII áttekintés, fő korlátok és lehetőségek, kezdeti kockázatok	részletes PMESII/METOC/geospatial elemzés frissítése és hiányok kijelölése	korlátozások/lehetőségek, kulcstényezők, indikátor-katalógus, gyűjtési terv
3. Adversary és aktorok értékelése	ellenség/aktork azonosítása, kezdeti szándék- és képességkép	képességek, szervezet, doktrína, logisztika, C2, hálózatok mély elemzése	actor profile-ok, erőrendi vázlat, hálótérképek, sebezhetőségi pontok
4. COA-k meghatározása	lehetséges ellenség szándékok gyors felvázolása, legvalószínűbb és legveszélyesebb irányok	SSA-hoz illesztett ellenséges COA-k, indikátorok és kiváltók, korai jelzőrendszer	COA-sablonok, eseményláncok, indikátorok és figyelési listák, valószínűségi értékelés

6-4. táblázat: JIPOE és COPD 1-2 fázisok kapcsolódása (saját szerkesztés) [33], [34], [10].

6.3. Adatgyűjtési módszerek és adatforrások: HUMINT, SIGINT, GEOINT/IMINT, OSINT, CYBINT

A COPD 1-2 fázisában az információs igények kijelölése akkor hatékony, ha a törzs a gyűjtési diszciplínák képességeit és korlátait ismeri, és a kérdéseket ennek megfelelően fogalmazza meg. A JDP 2-00 a hírszerzési folyamatot a gyűjtés, feldolgozás, elemzés és disszemináció rendszereként írja le, amelynek alkalmazkodónak és dinamikusnak kell lennie, miközben megőrzi a módszerességet [36].

Az alábbi alfejezetek röviden összefoglalják a fő adatgyűjtési módszereket és azt, hogy az MI mely pontokon tudja támogatni az előfeldolgozást és a minőségbiztosítást. A hangsúly a tervezői felhasználhatóságon van: mit várhat a törzs az adott forrástól a COPD 1-2 fázis időkeretében, és milyen kockázatokkal kell számolni.

6.3.1. HUMINT

A JDP 2-00 szerint a HUMINT olyan hírszerzési információ, amelyet emberi operátorok gyűjtenek, és amelyet elsősorban emberi források szolgáltatnak; megvalósulhat megfigyeléssel vagy közvetlen kommunikációval, és magában foglalja például a debriefinget, forráskezelést, taktikai kikérdezést, kihallgatást és kapcsolódó tevékenységeket [36].

A COPD 1-2 fázisában a HUMINT különösen értékes a szándék, a motiváció, a percepciók és a legitimációs mintázatok megértésében, továbbá ott, ahol a technikai felderítés korlátozott (pl. urbanizált környezet, civil-ruha és rejtett hálózatok).

MI-támogatás: a HUMINT anyagok (jelentések, jegyzetek, hangfelvételek átiratai) esetén a természetesnyelv-feldolgozás (NLP) segítheti a személy- és szervezetazonosítást, a kapcsolatok kinyerését, az esemény-idővonal összeállítását, illetve a kulcstémák gyors áttekintését. A HUMINT sajátossága azonban a forrásvédelem: az MI-feldolgozásnál különösen fontos a hozzáférés-szabályozás és az adatminimalizálás.

6.3.2. SIGINT

A JDP 2-00 a SIGINT-et az elektromágneses jelekből vagy emissziókból származó hírszerzésként írja le; ide tartozik a kommunikációs hírszerzés (COMINT) és az elektromágneses hírszerzés (ELINT) is [36]. A SIGINT a korai fázisban gyors képet adhat a kommunikációs aktivitásról, hálózati mintázatokról, illetve bizonyos platformok és rendszerek jelenlétéről.

MI-támogatás: jel- és metaadat szinten a gépi tanulás alkalmas lehet anomáliák, klaszterek és új aktivitási mintázatok detektálására (például új adó-azonosítók, frekvenciahasználat, forgalmi szint változása). A korai figyelmeztetéshez a trend- és változásdetektálás gyakran értékesebb, mint a teljes tartalmi megfejtés.

Korlátok: a SIGINT sok esetben erősen kontextusfüggő, és a műveleti környezetben a kibertér, a civil kommunikáció és a katonai rendszerek összefonódása miatt a téves attribúció kockázata nő. Ezért a COPD 1-2 fázisában a SIGINT megállapításokat célszerű más forrásokkal (GEOINT, HUMINT, OSINT) ellenőrizni.

6.3.3. GEOINT/IMINT

A JDP 2-00 szerint a GEOINT geospatial információ, képanyag és más adatok kiaknázásából és elemzéséből származó hírszerzés, amely földrajzilag referált tevékenységeket és jellemzőket ír le, értékel vagy vizuálisan ábrázol; a GEOINT magában foglalja az IMINT-et és alátámasztja a tervezést, navigációt és célkijelölést [36].

Ugyanebben a doktrínában az IMINT olyan hírszerzés, amely szenzorokkal (földi, tengeri, légi vagy űr platformok) beszerzett képanyagból származik; fontos megjegyzés, hogy a kép önmagában nem hírszerzés, az IMINT-et az elemzők által végzett értelmezés és értékelés hozza létre [36].

MI-támogatás: a számítógépes látás (CV) segíti a nagy mennyiségű kép és videó előszűrését (objektumdetektálás, változástetektálás, útvonal- és konvojfelismerés), valamint a tér-idő alapú anomáliák jelzését. Ez különösen hasznos a COPD 1-2 fázisában, ahol a gyors tájékozódás és a kritikus helyszínek azonosítása prioritás.

6.3.4. OSINT

A JDP 2-00 az OSINT-et nyilvánosan elérhető, illetve korlátozott hozzáférésű, de nem minősített információkból származó hírszerzésként határozza meg; a nyilvánosan elérhető információ (PAI) fogalma kiterjed az online tartalmakra, a publikált vagy sugárzott anyagokra, a nyilvános kéréssel hozzáférhető adatokra, valamint a paywall mögötti vagy kereskedelmi alapon gyűjtött információkra is [36].

A COPD 1-2 fázisában az OSINT azért kiemelt, mert gyorsan skálázható, és a politikai, társadalmi és információs dimenziókban gyakran ez adja a legkorábbi jelzéseket. Ugyanakkor az OSINT a leginkább manipulálható: dezinformáció, koordinált befolyásolás, hamis fiókok és szintetikus tartalom torzíthatja.

MI-támogatás: az NLP képes nagy szövegtömegből kulcstémákat, szereplőket és narratívákat kinyerni, a hálózatelemzés pedig a terjesztési mintázatokat, koordinált viselkedést és bot-hálózatokat azonosítani. A minőségbiztosítás része lehet a forrás-profilozás és a bizonytalanság jelölése is.

6.3.5. CYBINT (kibertéri információk és fenyegetési hírszerzés)

A kibertérből származó információk (CYBINT) a gyakorlatban hálózati forgalmi mintázatok, incidensjelentések, rosszindulatú infrastruktúrák (domain, IP), sérülékenységi információk, valamint fenyegető szereplők eszköztárának (TTP-k) elemzését jelentik. Bár a CYBINT elnevezés nem minden doktrínában jelenik meg önálló gyűjtési diszciplínaként, a JDP 2-00 részletesen tárgyalja az információs környezet kognitív, fizikai és virtuális dimenzióját, ami jó keretet ad a kibertéri jelenségek tervezői értelmezéséhez [36].

A COPD 1-2 fázisában a CYBINT két okból kritikus: (1) a válságok jelentős részében a kibertéri műveletek az első jelzések között jelennek meg (felderítés, zavarás, adatlopás, infrastruktúra elleni támadás); (2) a dezinformáció és a technikai támadások gyakran összehangoltan működnek (információs művelet és kiberművelet együtt).

MI-támogatás: anomáliaészlelés a hálózati forgalomban, klaszterezés az eseménylogokban, gépi tanulás a rosszindulatú minták felismerésére, valamint tudásgráfok a TTP-k és szereplők összekapcsolására. A CYBINT esetén a pontosság mellett a gyorsaság is fontos: a korai jelzés sokszor még bizonytalan, ezért a valószínűségi értékelés és a forrás-bizonytalanság jelölése alapvető.

6.3.6. Összefoglaló: források, adatfajták és MI-támogatási pontok

A 6-5. táblázat összegzi a fő diszciplínákat, a tipikus adatfajtákat és azokat a pontokat, ahol MI-alkalmazások a legnagyobb hatást érhetik el a COPD 1-2 fázis időkorlátaival mellett.

Diszciplína	Adatfajta	Erősség a COPD 1-2-ben	MI-támogatási fókusz	Fő kockázat
HUMINT	jelentés, átirat, megfigyelés	szándék, motiváció, legitimáció	NLP: entitás, kapcsolat, összegzés	forrásvédelem; torzítás

SIGINT	jel, metaadat	aktivitási mintázatok, C2 jelek	anomália; klaszter; trend	attribúció; kontextus
GEOINT/IMINT	kép, téradat	helyzet és változás vizuális bizonyíték	CV: objektum/változás	álcázás; időjárás
OSINT	média, közösségi	gyors, széles lefedés	NLP + hálózat: narratíva, bot	dezinformáció; zaj
CYBINT	log, IOC, TTP	korai jelzés, infrastruktúra kockázat	anomália + gráf	hamis zászló; érzékeny adatok

6-5. táblázat: Össz-adatforrású adatgyűjtés és MI-támogatási pontok (saját szerkesztés) [36].

6.4. Adatfúzió és elemző lánc (intelligence cycle); dezinformáció és deception kezelése

A többforrású adatgyűjtés önmagában nem garantál jobb helyzetértést. A döntés-előkészítés valódi szűk keresztmetszete az adatfúzió és az elemző lánc: hogyan lesz a heterogén adatokból értelmezett információ, majd megbízható elemzői ítélet, amelyet a parancsnok és a politikai döntéshozó fel tud használni.

A JDP 2-00 a hírszerzési ciklust négy átfogó funkcióra bontja: direction, collection, processing, dissemination (DCPD). A hangsúly az egyidejűségen és az átfedéseken van: a követelmények irányt adnak a gyűjtésnek, a feldolgozás és értékelés folyamatosan visszacsatol, a termék pedig visszakerül a döntéshozóhoz új igények generálására [36].

A JP 2-01.3 kiemeli, hogy a J-2 állománynak folyamatosan értékelnie és frissítenie kell a JIPOE termékeit, és biztosítani kell, hogy azok időben, pontosan, használhatóan, teljesen, objektíven és relevánsan elégítsék ki a parancsnok igényeit [10]. A COPD 1-2 fázisaiban ez különösen fontos, mert a stratégiai helyzetértékelés részben a még hiányos és gyorsan változó információs bázisra épül.

Az MI a fúziós láncban leginkább a feldolgozás és az elemzői előkészítés területén ad kézzelfogható nyereséget: automatikus nyelvi feldolgozás, képi előszűrés, metaadat-normalizálás, duplikátum- és ellentmondás-kezelés, valamint a valós idejű jelzések (alerting)

támogatása. A magasabb rendű elemzői ítélet azonban továbbra is emberi felelősség, különösen ott, ahol a bizonytalanság magas vagy ahol a döntés stratégiai következményekkel jár.

6.4.1. Adatfúziós architektúra: a nyers adatoktól a döntéstámogató termékig

A gyakorlati megvalósításban a össz-adatforrású fúzió három rétegben írható le:

- (1) Adat- és jel-szint (data-level): szenzor- és logadatok, nyers képek, jelparaméterek, hálózati metaadatok. Itt a fő feladat a minőségellenőrzés, idő- és helynormalizálás, valamint a zaj csökkentése. MI-eszközök: anomáliaészlelés, szűrés, automatikus címkézés.
- (2) Információ- és esemény-szint (information/event-level): entitások és események összekapcsolása, duplikátumok feloldása, a többforrású megerősítés (corroboration) és ellentmondások jelölése. MI-eszközök: entitás-összerendelés, eseményfúzió, tudásgráf.
- (3) Tudás- és döntés-szint (knowledge/decision-level): hipotézisek, következtetések, valószínűségek, alternatív értelmezések és indikátorok. MI-eszközök: előrejelzés, szcenárió- és érzékenység-elemzés, hálózati központosság és hatásmodellek.

A fúzió célja, hogy a törzs a COPD 1. fázisában gyorsan azonosítsa a válság relevanciáját és a kritikus kérdéseket, a 2. fázisban pedig olyan stratégiai értékelést készítsen, amelyben a kulestényezők és aktorok, a kockázatok, valamint a lehetséges ellenséges cselekvési mintázatok összhangban vannak [33].

6.4.2. Dezinformáció és deception: fenyegetések és kezelési módszerek

A korai helyzetértékelés egyik legkritikusabb kockázata a szándékos információtorzítás. A dezinformáció célja a befogadó megtévesztése hamis vagy félrevezető információk létrehozásával és terjesztésével; a deception katonai értelemben ennek a műveleti, taktikai és technikai megfelelője, amikor a saját szándék elrejtése és az ellenség félrevezetése a cél. A modern információs környezetben mindkettő gyakran többdimenziós: technikai, narratív és pszichológiai elemeket is tartalmaz.

A dezinformáció és a mesterséges intelligencia metszetében a generatív MI képes hiperrealisztikus szintetikus média (szöveg, hang, videó) előállítására, a terjesztés automatizálására (botok, hamis fiókok), valamint célzott, személyre szabott üzenetek

skalázására. Ugyanakkor a szakirodalom hangsúlyozza, hogy a hatások empirikus megítélése vegyes, és árnyalt, bizonyítékokon alapuló megközelítésre van szükség [39].

A COPD 1-2 fázisában a kezelési cél kettős: (1) csökkenteni a hamis információk bejutását a döntéstámogató termékekbe, (2) azonosítani a deception jeleit, és a bizonytalanságot transzparensszen kommunikálni. Ebben a folyamatban a hírszerzési elemzés minőségi standardjai kulcsfontosságúak: az Analytic Standards (ICD 203) többek között előírja a források és információk minőségének mérlegelését, az alternatív magyarázatok vizsgálatát, valamint a bizonytalanság és az indikátorok jelölését [37].

Fenomenológia	Tipikus jel (indikátor)	Kockázat a COPD 1-2-ben	Ellencsapás (tradecraft + MI)
Szintetikus tartalom	hang/videó anomáliák; újonnan létrehozott csatornák; gyors virális terjedés	hamis szándék- és képességkép	tartalom-hitelesítés; forrás-profil; multimodális detektálás
Koordinált terjesztés	fiók-koreográfia; azonos szövegvariánsok; időzített posztolás	narratíva torzítása; percepciók manipulálása	hálózat-elemzés; bot-detektálás; platform-jelentések
Deception a fizikai térben	álcázás; dummy eszközök; szándékos zaj	téves erőrendi következtetés	többforrású megerősítés; CV + változásdetektálás; HUMINT ellenőrzés
Kibertéri hamis zászló	IOC-k keverése; új infrastruktúra; indikátorok ellentmondása	hibás attribúció; rossz kockázatértékelés	TTP-alapú elemzés; tudásgráf; bizonytalanság jelölése
Adatminőségrombolás	adatforrás hirtelen minőségromlása; metaadat hiány; manipulált időbélyeg	hibás trend és predikció	adatproveniencia; validáció; anomália a metaadatban

6-6. táblázat: Dezinformáció és deception kezelése indikátorokkal és MI-támogatással (saját szerkesztés) [37], [39].

6.4.3. MI-kockázatok és kontrollok az elemző láncban

Az MI bevezetése önmagában új kockázatokat is hoz: modell-hallucináció, torzítás (bias), robusztussági és biztonsági problémák, illetve a támadó fél által kiváltott hibák (adversarial

inputok). Ezek a kockázatok a COPD korai fázisaiban azért különösen veszélyesek, mert a döntési ablak rövid, és a hibák gyorsan beépülhetnek a stratégiai értékelésbe.

A NIST AI Risk Management Framework (AI RMF) a megbízható MI dimenziói között kiemeli a valid és megbízható működést, a biztonságot, a kiberbiztonsági rezilienciát, az elszámoltathatóságot és az átláthatóságot, valamint a magyarázhatóságot és értelmezhetőséget [12]. A hírszerzési és tervezési környezetben a gyakorlati követelmény az, hogy az MI által adott jelzésekhez legyen visszakereshető magyarázat, legyen dokumentált adatproveniencia, és legyen emberi felülvizsgálati pont (human-on-the-loop) a kritikus termékekbe történő beemelés előtt.

Ennek megfelelően javasolt kontrollok: (1) adatminőség és provenance követelmények (forrás és feldolgozási lépések naplózása), (2) modellvalidáció és drift-monitorozás, (3) red teaming és adversarial tesztelés, (4) hozzáférés-kezelés és információbiztonság, (5) a bizonytalanság explicit jelölése és a döntéstámogató termékben történő kommunikálása.

6.5. Fenygetésértékelés (Threat Assessment) prediktív elemzése

A COPD 1-2 fázisában a fenygetésértékelés célja nem a teljes bizonyosság, hanem a döntéshez releváns valószínűségi állítások és kockázati következmények megfogalmazása. A JIPOE 4. lépése (COA-k meghatározása) olyan termékeket ad, amelyek az ellenség lehetséges cselekvési irányait, a kiváltó eseményeket és az indikátorokat rendszerezik [10].

A prediktív elemzés itt tág értelemben szerepel: magában foglalja az időbeli trendek értékelését, az indikátor-alapú korai jelzést (early warning), a valószínűségi forgatókönyv-elemzést, valamint a hálózati és tér-idő dinamikákból levezetett előrejelzést. A modern hírszerzési elemzésről szóló szakirodalom arra hívja fel a figyelmet, hogy a technológia akkor teremti tartós értékét, ha az elemzői módszertan, a folyamat és az emberi szerepek együtt fejlődnek [40].

A prediktív eszközök alkalmazása a COPD 1-2 fázisában három gyakorlati elvhez köthető: (1) indikátorok és küszöbértékek explicit kijelölése; (2) bizonytalanság és alternatív magyarázatok jelölése; (3) folyamatos frissítés új adatokkal. Az ICD 203 ösztönzi, hogy az elemzés jelezze a bizonytalanság okait és azokat az indikátorokat, amelyek a fő megállapításokat megváltoztathatják [37].

6.5.1. Prediktív módszerek a fenyegetésvértékelésben

A 6-7. táblázat áttekintést ad a tervezési célú fenyegetésvértékelésben alkalmazható prediktív módszerekről, az elvárt inputokról és az eredmények értelmezési kockázatairól.

Módszer	Bemenet	Tipikus kimenet	Erősség	Korlát / kockázat
Indikátor-alapú korai jelzés	kiváltók és mérőszámok (PMESII)	riasztás, trendjel	átlátható, gyors	rossz indikátor-választás; deception
Idősor és változásdetektálás	időbélyeges események	szintváltás, szezonális hatás	jó a gyors változásokra	adatminőség; drift
Gráf- és hálózati predikció	szereplő-kapcsolat háló	kulcsszereplők, terjedési pályák	rendszer szemlélet	hiányos háló; attribúció
Tér-idő modellek	georeferált események	hotspot, útvonal	helyzetképhez illeszkedik	rejtett változók
Szimuláció és wargaming	szabályok, feltételezések	forgatókönyvek	mi lenne, ha elemzés	modell-feltevés érzékeny
ML osztályozás/regresszió	többforrású jellemzők	valószínűség/score	skálázható	magyarázhatóság; túlillesztés

6-7. táblázat: Prediktív módszerek fenyegetésvértékeléshez (saját szerkesztés) [10], [37], [12], [40].

6.5.2. Minőségbiztosítás és ember-gép együttműködés

A prediktív elemzés erősen érzékeny az adatok torzítására és hiányosságaira. Ezért a COPD korai fázisaiban a predikciót célszerű kockázati hipotézisként kezelni: a cél, hogy a törzs kijelölje, mely információk csökkentenék a bizonytalanságot, és mely indikátorok esetén kell a stratégiai megállapításokat frissíteni.

A gyakorlatban ez ember-gép együttműködés. Az MI gyorsan számol, mintázatokat jelez és rangsorol, az elemző pedig kontextusba helyezi a jelzéseket, alternatív magyarázatokat vizsgál, és dönt a felhasználhatóságról. A NIST AI RMF megközelítése alapján a validáció, a monitoring és a kockázatkezelés életciklus-szintű feladat, nem egyszeri bevezetési lépés [12].

6.6. MI eszközök az elemzés támogatására (NLP, CV, gráf- és hálózatelemzés, anomáliaészlelés)

A COPD 1-2 fázisában az MI alkalmazásának célja a döntéstámogató termékek gyorsabb, konzisztensebb és transzparensabb előállítása. A hangsúly nem egyetlen „nagy rendszer” kialakításán, hanem moduláris képességek beillesztésén van: olyan eszközökön, amelyek a JIPOE és az intelligence cycle egyes pontjain automatikusan előkészítik az adatot, jelzik a változást, és csökkentik az elemzők terhelését [36], [10].

A bevezetésnél alapvető követelmény, hogy az MI kockázatai kezelhetők legyenek. A NIST AI RMF a megbízhatóságot és a kockázatkezelést a Govern–Map–Measure–Manage logikában tárgyalja, és a katonai-tervezési felhasználásban is jól alkalmazható a kontrollpontok kijelölésére [12].

A következőkben a fő MI-eszközcsoportokat a tervezői feladatokhoz kötve mutatjuk be: milyen bemenetekkel dolgoznak, milyen kimeneteket adnak, és hol illeszthetők be a COPD 1-2 fázis munkaritmusába.

6.6.1. Természetesnyelv-feldolgozás (NLP)

Az NLP eszközök az elemző lánc leggyakoribb bemenetét kezelik: szöveges jelentések, hírek, közösségi média, hivatalos közlemények, átiratok és dokumentumok. A COPD 1-2 fázisában az NLP tipikusan az alábbi feladatokban ad gyors nyereséget:

- Entitás- és kapcsolatkinyerés: személyek, szervezetek, helyek, eszközök, események azonosítása és összekapcsolása a közös adatmodellben.
- Esemény- és idővonalépítés: események időbeli sorrendbe rendezése, kiváltók és következmények jelölése.
- Téma- és narratívaelemzés: domináns témák, visszatérő keretezések, propaganda- és dezinformációs narratívák feltérképezése.
- Többnyelvű támogatás: gépi fordítás és terminológiai normalizálás, különösen többnemzeti törzskörnyezetben.
- Összefoglalás és priorizálás: hosszú dokumentumok rövid kivonata, releváns részletek kiemelése az elemzői követelmények alapján.

E képességek közvetlenül segítik a JIPOE 1-2. lépését (OE definiálása és hatásainak leírása), mert a szereplők és dinamikák gyorsabb azonosításával csökkentik a kezdeti bizonytalanságot.

A dezinformáció ellen a forrás- és terjesztési mintázatok vizsgálata, valamint a tartalom állítás-szintű elemzése (claim extraction) adhat alapot [37], [39].

Korlátok és kontrollok: a generatív nyelvi modellek esetén kritikus a hallucinációk kezelése, a források explicit hivatkozása és a magyarázhatóság. Tervezői környezetben ajánlott az MI-t olyan feladatokra használni, ahol a hibahatás kontrollálható (előszűrés, jelzés, összegzés), míg a végső megállapításokat az ICD 203 szerinti elemzői standardok szerint kell megfogalmazni [37].

6.6.2. Számítógépes látás (Computer Vision, CV) és térképi analitika

A JDP 2-00 megjegyzi, hogy a mesterséges intelligencia és a gépi tanulás képes támogatni a képanyag elemzésének bizonyos aspektusait, miközben az IMINT továbbra is elemzői értelmezést igényel [36]. A COPD 1-2 fázisában a CV-eszközök főleg a feldolgozás gyorsításában és a változás jelzésében adnak értéket.

Tipikus alkalmazások:

- Objektumdetektálás és osztályozás: járművek, eszközök, fegyverrendszerek, építmények azonosítása.
- Változásdetektálás: táborok, raktárak, útvonalak, ellenőrzőpontok megjelenése vagy bővülése; infrastruktúra-sérülések.
- Mozgás- és mintázatelemzés: konvojok, hajóforgalom, repülőterek aktivitása, időbeli ritmusok.
- Georeferált vizualizáció: a CV által detektált elemek térképi rétegekbe szervezése, ami a közös helyzetkép (COP) gyors építését támogatja.

Korlátok: az álcázás, a rossz minőségű szenzoradat, a kedvezőtlen meteorológiai viszonyok és a szándékos deception mind növelhetik a téves észlelés esélyét. Ezért a CV-eredményeket célszerű többforrású megerősítéssel (SIGINT/HUMINT/OSINT) és emberi ellenőrzéssel használni, különösen akkor, ha a következtetés stratégiai szintű döntést befolyásol.

6.6.3. Gráf- és hálózatelemzés

A PMESII/ASCOPE és az ICR2 jellegű keretek természetesen hálózatos reprezentáció felé terelik az elemzést: szereplők, erőforrások, kommunikációs csatornák és szabályok kapcsolatrendszere határozza meg, hogy egy beavatkozás hol és hogyan fejti ki hatását. A gráf- és hálózatelemzés ezért a COPD 1-2 fázisában különösen hasznos a komplex, hibrid fenyegetések értelmezésében.

A gráfmodell előnye, hogy ugyanabban a struktúrában ábrázolhatóak a katonai, politikai, gazdasági és információs kapcsolatok: például egy szereplő pénzügyi háttere, kommunikációs hálózata és fegyveres alárendeltségi rendszere. Az all-source szemléletű hírszerzési ökoszisztéma éppen azt igényli, hogy a különböző forrásokból származó mozaikdarabok konzisztens képpé álljanak össze [41].

Tipikus alkalmazások:

- Link analysis és központisági mutatók: kulcsszereplők, közvetítők, kritikus csomópontok azonosítása.
- Közösségetektálás: frakciók, sejtek, érdekhálózatok és ellátási alhálózatok feltérképezése.
- Terjedési modellek: narratívák, radikalizációs mintázatok, dezinformációs kampányok diffúziója.
- Útvonal- és ellátási lánc elemzés: logisztikai függőségek, kritikus útvonalak és szűk keresztmetszetek.

MI-támogatás: automatikus gráfépítés (NLP-ből és metaadatból), gráf-beágyazások és prediktív link-elemzés. A kimenetek azonban csak akkor használhatók döntéstámogatásra, ha a forrásbizonytalanság és a hiányos hálózatból eredő torzítások explicit jelölést kapnak az ICD 203 elvei szerint [37].

6.6.4. Anomáliaészlelés és riasztás (alerting)

Az anomáliaészlelés a COPD 1-2 fázisában különösen értékes, mert a válságok kezdeti jelei gyakran nem egyetlen „nagy eseményben”, hanem sok kicsi, statisztikailag szokatlan eltérésben jelennek meg. Ilyen lehet például egy kommunikációs mintázat megváltozása (SIGINT), egy infrastruktúra üzemállapotának rendellenessége (GEOINT/OSINT), vagy egy új rosszindulatú hálózati viselkedés (CYBINT).

Az MI-alapú anomáliaészlelés tipikusan nem „igaz/hamis” állítást ad, hanem prioritási listát: mely jelzések igényelnek emberi ellenőrzést. A döntéstámogatási érték akkor nő, ha az anomáliák összevethetők az előre definiált indikátorokkal, és ha a riasztás magyarázható (mely jellemzők tértek el, milyen erős a jel, milyen alternatív magyarázatok lehetségesek) [12].

Az anomáliaészlelésnek van egy másik, kevésbé látványos, de kritikus szerepe: adatminőség-védelem. A hirtelen adatminőség-romlás, metaadat-manipuláció vagy forrás-csatorna elnémulása is anomália, amely a fúziós lánc egészét torzíthatja. A korai fázisban ezért célszerű külön „adatminőség indikátorokat” fenntartani a legfontosabb forrásokra.

6.6.5. Integráció a törzsmunkába és termékekbe (COP, SSA, briefek)

Az MI-eszközök akkor hasznosak a COPD 1-2 fázisában, ha a kimenetek közvetlenül illeszkednek a törzs döntéstámogató termékeihez: térképi rétegek, trendgrafikonok, indikátor-listák, actor profile-ok, COA-sablonok és rövid vezetői összefoglalók. A JP 2-01 a hírszerzési támogatást a műveletek egészéhez kapcsolja, és hangsúlyozza a releváns információk időbeni eljuttatásának fontosságát a parancsnoki döntéshozatal támogatására [38].

A JP 5-0 a tervezést olyan iteratív folyamatként kezeli, amelyben a helyzetértés, a feltételezések és a kockázatok folyamatosan frissülnek. A COPD 1-2 fázisában ez a frissítési ciklus különösen gyors: az MI-alapú feldolgozás értéke akkor maximalizálható, ha a frissítések automatikusan visszajutnak a JIPOE-termékekbe és a vezetői briefekbe [42].

Gyakorlati architektúra (ajánlás):

- Adatbeáramlás: ISR diszciplínák, OSINT, kibertéri logok, diplomáciai és partnerségi jelentések.
- Adatkezelés: metaadat-standardizálás, provenance napló, hozzáférés-kezelés (need-to-know).
- Tudásréteg: közös entitás- és eseményszótár, tudásgráf a szereplők és kapcsolatok kezelésére.
- Elemzői munkafelület: keresés, idősor, térkép, háló, és automatikus jelzések (alerting).
- Termékcsatorna: COP rétegek, SSA-fejezetek, COA-kiegészítések, indikátor és kockázati briefek.
- Visszacsatolás: elemzői felülvizsgálat, modell-monitorozás, új információigények generálása.

Az architektúra célja, hogy az MI ne párhuzamos világot hozzon létre, hanem a törzs meglévő folyamatait gyorsítsa és támogassa. A kritikus pont a hitelesség és a transzparencia: minden automatikus jelzés mellett láthatóvá kell tenni a forrásokat, a feldolgozási lépéseket és a bizonytalanságot, különösen a dezinformációval terhelt információs környezetben [37], [12], [39].

6.7. Részletes módszertani kiegészítések és megvalósítási javaslatok

Az alábbi alfejezetek a 6.1–6.6 pontok gyakorlati alkalmazását mélyítik el. A cél az, hogy a COPD 1–2 fázisában a törzs számára kézzelfogható sablonokat, ellenőrző listákat és beilleszthető MI-funkciókat adjunk. A kiegészítések a multi-source működésből indulnak ki, és külön hangsúlyt helyeznek a minőségbiztosításra, mert a dezinformációval terhelt információs környezetben a gyorsaság önmagában nem elég [37], [39].

6.7.1. Indikátor-katalógus kialakítása (PMESII-PT/ASCOPE szemlélet)

Az indikátorok akkor hasznosak, ha nem pusztán lista, hanem egy logikusan felépített katalógus részei. A PMESII-PT és az ASCOPE kombinációja segít abban, hogy az indikátorok ne csak absztrakt dimenziókhoz, hanem konkrét területekhez, struktúrákhoz és szereplőkhöz kapcsolódjanak [35].

Javasolt indikátor-sablon (minimum mezők): (1) név és leírás, (2) PMESII-PT és ASCOPE hozzárendelés, (3) megfigyelhető jelenség és mérési módszer, (4) primer és alternatív adatforrások, (5) küszöbérték és riasztási logika, (6) lehetséges magyarázatok és deception-rizikó, (7) várható következmények (mit változtat a SSA-ban), (8) felelős elemző és felülvizsgálati ciklus.

A katalógus kialakításakor célszerű a „döntési relevancia” elvét követni: minden indikátornak legyen explicit kapcsolata legalább egy döntési kérdéshez (pl. „Milyen valószínű, hogy az ellenség A irányba eszkalál?”, „Mekkora a civil kockázat egy adott műveleti opció mellett?”). Ez összhangban van az ICD 203 elvével, amely szerint az elemzésnek jelölnie kell azokat az indikátorokat, amelyek a fő megállapításokat megváltoztathatják [37].

6-8. táblázat egy példán keresztül mutatja be, hogyan nézhet ki egy indikátor-sablon kitöltve. A példa illusztratív, célja a módszertan bemutatása.

Mező	Példa kitöltés
Indikátor neve	Kritikus infrastruktúra zavarásának kockázata (elektromos hálózat)
PMESII-PT	Infrastruktúra + Fizikai környezet + Idő
ASCOPE	Structures (alállomások), Capabilities (helyreállítás), Events (kimaradások)
Megfigyelhető jel	gyakoribb rövid kimaradások; alállomás környéki aktivitás; kibertéri bejutási kísérletek
Primer források	GEOINT/IMINT (fizikai aktivitás), CYBINT (IOC, log), OSINT (üzemeltetői közlemény)
Küszöb/riasztás	a rövid kimaradások gyakorisága 30 napos átlag felett + új IOC megjelenése az ágazati CERT-ben
Alternatív magyarázat	időjárás; karbantartás; véletlen hiba; szabotázs; deception
Döntési következmény	SSA kockázati fejezet frissítése; erőforrás-átcsoportosítás CIP-re; kommunikáció előkészítése
Felülvizsgálat	napi monitorozás; heti kalibráció; esemény után post-mortem

6-8. táblázat: Indikátor-sablon példa (saját szerkesztés) [35], [37].

6.7.2. JIPOE termékek részletesen és MI-vel támogatható előállításuk

A JP 2-01.3 a JIPOE-t nem csak lépéssorozatként, hanem termékcsaládként is kezeli: a térképi és analitikai termékek célja, hogy a parancsnok és a törzs közös nyelven beszéljen a környezetről és az ellenség lehetséges cselekvési irányairól [10]. A COPD 1-2 fázisában ezek a termékek gyakran „minimum viable” formában készülnek el, majd iteratívan bővülnek.

Tipikus JIPOE termékek (tervezői nézőpontból):

- OE-leíró összefoglaló és alaptérképek (AO, AOI, releváns régiók).
- MCOO és korlátozó tényezők (terep, időjárás, infrastruktúra, populáció).
- Actor profile-ok (képesség, szándék, szervezet, doktrína, logisztika, C2).
- Situation template: a várható ellenséges település és tevékenységek vázlata.
- Event template és event matrix: eseménysorok, kiváltók, indikátorok és megfigyelési pontok.

- COA-sablonok: legvalószínűbb és legveszélyesebb ellenség-COA, hozzájuk rendelt figyelési lista.

MI-támogatási lehetőségek a termékekben:

- CV-alapú változásdetektálás a MCOO és infrastruktúra-korlátok frissítéséhez.
- NLP-alapú actor profile építés (szereplők céljai, retorika, kapcsolatok) HUMINT/OSINT anyagokból.
- Gráf- és hálózatelemzés a befolyási hálók és logisztikai függőségek feltérképezésére.
- Anomália és riasztás a kiváltók és indikátorok automatikus figyeléséhez.

A JP 2-01.3 hangsúlyozza, hogy a JIPOE termékeket folyamatosan értékelni és frissíteni kell, és biztosítani kell, hogy azok a tervezés tempójához igazodjanak [10]. Ez a követelmény a gyakorlatban MLOps-szerű működést jelent: adatfrissítés, modell-monitorozás, emberi ellenőrzés és dokumentált változáskezelés.

6.7.3. Gyűjtési követelmények és collection management a COPD 1-2 fázisában

A COPD 1-2 fázisában a gyűjtési feladatok akkor adnak gyors eredményt, ha a törzs a kérdéseket a döntési problémákból vezeti le. A COPD 2. fázisa során a stratégiai környezet elemzése közben kifejezetten azonosítani kell a hiányzó információkat és tudást, és ezeket a hiányokat a megfelelő szereplőkhöz kell rendelni [33].

A JDP 2-00 szerint a direction funkció a hírszerzési ciklus alapja: a parancsnoki döntési igényekből követelmények keletkeznek, amelyek irányítják a gyűjtést, és amelyek később visszacsatolás alapján módosulnak [36]. A JP 2-01 a gyűjtés és a disszemináció összehangolását hangsúlyozza, hogy a releváns információk időben érkezzenek a döntéshozókhöz [38].

Javasolt gyakorlati lépések a korai fázisban:

- Döntési kérdések bontása információigényekre: mit kell tudnunk ahhoz, hogy a SSA fő megállapításai megalapozottak legyenek.
- Információigények bontása indikátorokra: mit figyelünk, mi számít változásnak.
- Indikátorok bontása gyűjtési feladatokra: mely forrás, milyen formátum, milyen gyakoriság.

- Gyűjtési feladatok prioritizálása: időkritikusság, hatás, költség, kockázat és jogi korlátok alapján.
- Eredmények visszacsatolása: a beérkező információk alapján a feltételezések és a gyűjtési terv frissítése.

Az MI itt „követelmény-asszisztensként” jelenhet meg: a korábbi válságokból és jelentésekből tanult mintázatok alapján javasolhat indikátorokat, összekapcsolhat hasonló információigényeket, és automatikusan csoportosíthatja az RFIs-t a PMESII/ASCOPE logika szerint. A döntés azonban emberi: az indikátorok és a gyűjtési feladatok kiválasztása a parancsnoki szándék és a kockázati tolerancia függvénye [33], [37].

6.7.4. Elemzői tradecraft: strukturált technikák, torzítások és megbízhatóság

A dezinformációval terhelt környezetben a legjobb MI-eszközök sem helyettesítik az elemzői tradecraft-ot. Az ICD 203 analitikai standardjai olyan minimum követelményeket adnak, amelyek a COPD 1-2 fázisában is alkalmazhatók: források és információk minőségének mérlegelése, logikai következetesség, alternatív magyarázatok vizsgálata, bizonytalanság transzparens jelölése és időbeni aktualitás [37].

A modern hírszerzési elemzésről szóló irodalom felhívja a figyelmet arra, hogy a technológia nem csak gyorsít, hanem új torzításokat is hozhat (például túlzott bizalom a modell pontszámaiban, „automation bias”). Ezért az MI bevezetésével párhuzamosan a folyamatokat és a képzést is fejleszteni kell [40].

Javasolt strukturált technikák a korai fázisban (rövid ciklusokra adaptálva):

- Kulcsfeltevések (Key Assumptions Check): a SSA kritikus feltevéseinek explicit listázása és rendszeres felülvizsgálata.
- Versengő hipotézisek (ACH): az ellentmondó adatok strukturált kezelése, különösen deception gyanú esetén.
- Piros csapat (red teaming): ellenség szemszögű ellenőrzés, legveszélyesebb COA-k tesztelése.
- Indikátor-ellenőrzés: mi lenne az a jel, amely megcáfolná a fő állítást, és hogyan figyeljük ezt.

Az MI ezekben a technikákban támogató szerepben jelenhet meg: gyors keresés, releváns források összegyűjtése, ellentmondások jelzése, és a hipotézisekhez tartozó evidenciák rendszerezése. A végső ítélet, a megbízhatósági szint és a bizonytalanság kommunikálása azonban elemzői felelősség [37].

6.7.5. Prediktív modellek validációja és mérőszámok

A fenyegetésvértékelési predikciók értéke csak akkor fenntartható, ha mérhető. A katonai tervezésben gyakori probléma, hogy a prediktív állítások nem kerülnek utólagos ellenőrzésre, így a szervezet nem tanul a saját tévedéseiből. A NIST AI RMF hangsúlyozza a mérés (Measure) és a kezelés (Manage) fontosságát, beleértve a teljesítmény, robusztusság és kockázat monitorozását [12].

Javasolt mérőszámok és eljárások a hírszerzési predikciókhoz:

- Kalibráció: a valószínűségi becslések mennyire felelnek meg a bekövetkezési arányoknak (pl. Brier score).
- Diszkrimináció: mennyire különbözteti meg a modell a magas és alacsony kockázatú eseteket (pl. ROC-AUC, PR-AUC).
- Hasznosság: a predikció mennyiben segítette a döntést (pl. időnyereség, erőforrás-fókusz javulása).
- Drift és stabilitás: változik-e a modell teljesítménye a környezet változásával, és milyen gyorsan romlik.
- Deception ellenállás: mennyire érzékeny a modell a szándékos input-manipulációra és a forrásminőség romlására.

A validáció eredményeit célszerű visszacsatolni a JIPOE indikátorlistájába és a gyűjtési tervbe. A cél nem a „tökéletes modell”, hanem a jobb döntési támogatás: a modell ott legyen erős, ahol gyors jelzés kell, és ott legyen kontrollált, ahol a hiba stratégiai kockázatot jelent.

6.7.6. Bevezetési útiterv és szervezeti feltételek

A COPD 1-2 fázisában alkalmazott MI-eszközök bevezetése nem kizárólag technikai feladat. Szervezeti, jogi, információbiztonsági és képzési dimenziói is vannak, különösen többnemzeti környezetben. A nemzetbiztonsági és védelmi szakirodalom kiemeli, hogy a katonai MI

alkalmazások sikere a megfelelő adatinfrastruktúrán és a humán–gép együttműködésen múlik [43], [44]. A NATO mesterséges intelligencia stratégiája hangsúlyozza a felelős alkalmazást és az interoperabilitási követelményeket, amelyek a többnemzeti törzskörnyezetben különösen relevánsak [45].

Javasolt lépcsőzetes bevezetés:

- 1. lépcső: „low-risk” automatizálás (fordítás, keresés, duplikátum-kezelés, dokumentum-összegzés) szigorú forrásjelöléssel.
- 2. lépcső: előfeldolgozó MI a multi-source láncban (CV előszűrés, entitás-összerendelés, anomália-riadók), emberi validációval.
- 3. lépcső: döntéstámogató analitika (indikátor-dashboard, hálózatelemzés, prediktív score-ok), kontrollált felhasználással és mérőszámokkal.
- 4. lépcső: integrált munkafolyamatok (MLOps, drift-monitorozás, red teaming, post-mortem), szervezeti tanulással.

Adat- és információbiztonság: a HUMINT és minősített adatkezelés miatt a zero-trust jellegű hozzáférés-kezelés, a naplózás és a provenance dokumentálása alapkövetelmény. Az MI-rendszernek a NIST AI RMF szerint is kezelnie kell a biztonsági és robusztussági kockázatokat, beleértve a támadó fél általi befolyásolást [12].

Képzés és kultúra: a törzsnek meg kell tanulnia „MI-vel együtt dolgozni”. Ez egyszerre jelent technikai készséget (eszközhasználat) és elemzői fegyelmet (a bizonytalanság kommunikálása, alternatívák vizsgálata). A döntéstámogató rendszerekre vonatkozó kutatások is azt jelzik, hogy az MI előnye akkor realizálódik, ha a felhasználó érti a rendszer korlátait, és a folyamatok is alkalmazkodnak [46], [47].

6.9. Összegzés

A COPD 1-2 fázisai a stratégiai döntés-előkészítés legidőérzékenyebb részei. A siker kulcsa a műveleti környezet gyors, mégis módszeres értelmezése, a kritikus információhiányok kijelölése és a fenyegetésvértékelés valószínűségi megfogalmazása [33], [10].

A fejezet bemutatta, hogy a DIME/PMESII és kapcsolódó keretek (PMESII-PT, ASCOPE, ICR2) egy közös adatmodellben kezelhetők, ahol entitások, kapcsolatok, események és indikátorok adnak alapot a össz-adatforrású fűziónak. A JIPOE négy lépése a COPD 1-2

fázisaiban természetes „gerincfolyamatot” ad, amelybe a HUMINT, SIGINT, GEOINT/IMINT, OSINT és CYBINT adatai integrálhatók [34], [36], [10].

Az MI legnagyobb értéke a feldolgozás és előkészítés gyorsítása: NLP, számítógépes látás, hálózatelemzés és anomáliaészlelés képes csökkenteni a zajt, jelzéseket adni és strukturált adatot előállítani. Ugyanakkor a dezinformáció és deception környezetében a tradecraft és a kockázatkezelés elsődleges: az ICD 203 elemzői standardjai és a NIST AI RMF kontrollrendszere együtt teremthetnek olyan keretet, amelyben az MI haszna érvényesül, miközben a kockázatok kezelhetők [37], [12], [39].

A következő fejezetek számára természetes folytatás a 3-5. COPD fázisok MI-támogatásának tárgyalása (COA-fejlesztés, tervkidolgozás, wargaming és assessment), ahol a prediktív és szimulációs képességek, valamint a humán-gép együttműködés kérdései még hangsúlyosabbá válnak [42].

7. FEJEZET MI TÁMOGATÁS A COPD 3-4 FÁZISÁBAN: KATONAI VÁLASZ OPCIOK, KONCEPCIÓ ÉS TERVEZÉS (MILITARY RESPONSE OPTIONS, CONOPS, OPLAN)

A COPD (Comprehensive Operations Planning Directive) 3–4. fázisai a tervezési folyamat „döntési” és „kivitelezhetővé tételi” szakaszai: itt alakul át a műveleti környezet megértése és a stratégiai helyzetértékelés (1–2. fázis) olyan cselekvési alternatívákká, amelyekről a parancsnok/vezető döntést hozhat, majd olyan tervdokumentumokká, amelyek alapján a végrehajtás megszervezhető és irányítható. A COPD 3. fázisában a törzs több katonai válasz opciót (Military Response Options – MRO) és cselekvési tervet (Course of Action – COA) dolgoz ki, elemzi és hasonlítja össze. A folyamat célja, hogy a parancsnok számára egyértelműen bemutatható legyen: melyik COA hogyan éri el a kívánt végállapotot, milyen kockázatokkal és erőforrás-igénnyel jár, és milyen másod- vagy harmadrendű hatások várhatók. A COPD 4. fázisa ezt követően a kiválasztott COA-t koncepcióvá (CONOPS), majd részletes hadműveleti tervvé (OPLAN) formálja, mellékletekkel, függelékekkel és végrehajtási szabályokkal, összhangban a műveleti szintű tervezési alapelvekkel. [33], [34], [42]

A 3–4. fázis sajátossága, hogy a döntés-előkészítésnek egyszerre kell gyorsnak, következetesnek és auditálhatónak lennie. A COA-k tere a gyakorlatban nagy: a célok (ends, objectives), a rendelkezésre álló eszközök (means), az alkalmazás módja (ways), a politikai/jogi korlátok, a kockázatvállalási szint, valamint a koalíciós és civil szereplőkkel való koordináció együttesen kombinatorikus problémateret hoz létre. Ebben a térben a törzs egyszerre kényszerül kreatív tervezésre (új kombinációk, újszerű megoldások) és fegyelmezett elemzésre (korlátok, logisztika, idő–tér összhang, műveleti kockázat). AJP-5 és a JP 5-0 egyaránt hangsúlyozza, hogy a COA-elemzés és a háborús játék nem egyszeri „teszt”, hanem iteratív tanulási mechanizmus, amely a visszacsatolás révén folyamatosan finomítja a koncepciót és a részletes tervet. [34], [42]

A mesterséges intelligencia (MI) ebben a szakaszban elsősorban a „tervezési tér” strukturálásában, a változatok gyors generálásában, az erőforrás-elosztás optimalizálásában, valamint a háborús játék és szimulációk eredményeinek feldolgozásában tölthet be képességnövelő szerepet. Fontos hangsúlyozni: a MI nem „dönt” a parancsnok helyett. A felelősség, a legitimáció és a kockázatvállalás továbbra is emberi, ugyanakkor a MI képes olyan alternatívákat, érzékenységvizsgálatokat és következményláncokat feltárni, amelyek manuális

eljárásokkal nehezen láthatók át. A fejezet MI-megközelítése ezért az ember–gép együttműködésre (human-in-the-loop), a transzparens adat- és modellkormányzásra, valamint a kockázatkezelésre épül, összhangban a NIST AI RMF 1.0 és a NATO felelős MI-használatra vonatkozó irányainak logikájával. [12], [44], [45]

A fejezet a COPD 3–4. fázisain belül három gyakorlati fókusz tárgyal. (1) Az MI által vezérelt COA/MRO generálás és elemzés módszereit, beleértve a szimulációs és többforgatókönyves (Monte Carlo) értékelést. (2) Az erőforrás-elosztás optimalizálására alkalmazható algoritmusokat, azok előnyeit és korlátait. (3) Az MI-integráció lehetőségeit a háborús játékok (wargaming) és modellezés–szimuláció (M&S) eszközeivel.

7.1. Mesterséges intelligencia által vezérelt cselekvési terv (MRO/COA) generálása és szimulációja

A COPD 3. fázisában a COA-k kidolgozása és elemzése tipikusan három egymásra épülő lépésként jelenik meg: (1) COA-k generálása és kidolgozása (COA development), (2) COA-elemzés – gyakran háborús játék (COA analysis / wargaming), majd (3) COA-összehasonlítás és ajánlás a döntéshozónak (COA comparison). A JP 5-0 és a NATO tervezési logika szerint a „jó” COA-nak alkalmasnak (suitability), megvalósíthatónak (feasibility), elfogadhatónak (acceptability), megkülönböztethetőnek (distinguishability) és kellően kidolgozott/teljesnek (completeness) kell lennie. [33], [34], [42]

MI-támogatás szempontjából a COA-készítés két problémára vezethető vissza. Az első a keresési probléma: milyen „műveleti koncepció” variánsok felelnek meg a céloknak és korlátoknak? A második az értékelési probléma: a jelöltek közül melyik adja a legjobb kompromisszumot a hatásosság, a kockázat, a költség és az idő dimenzióiban, figyelembe véve a bizonytalanságot és az ellenség adaptív reakcióját. A MI mindkét ponton segíthet, de eltérő eszköztárral: a generálásnál strukturált tudásreprezentáció és (gépi) tervezés, az értékelésnél szimuláció, prediktív analitika és többkritériumos döntéstámogatás dominál.

A komplex, sokszereplős problémákra irányuló tervezési szakirodalom rámutat arra, hogy az operatív tervezés akkor hatékony, ha a törzs képes a katonai, politikai és civil szempontokat egyetlen koherens kampánylogikába integrálni, és a tervezés során folyamatosan tanulni a visszacsatolásokból. Ez a szemlélet jól illeszkedik a MI-támogatott tervezéshez, mert a MI a variánsok gyors generálásával és a visszacsatolások strukturálásával éppen ezt az iteratív tanulást erősítheti. [48]

7.1.1. Bemenetek, korlátok és tervezési „szerződés”: mit kell a MI-nek tudnia?

A COA-generálás nem „szövegírási feladat”, hanem formális és informális korlátokkal leírható döntési tér bejárása. A bemenetek nagy része már a COPD 1–2. fázisában létrejön (küldetés és felsőbb iránymutatás, JIPOE/J2 termékek, PMESII-PT alapú helyzetkép), de a 3. fázisban ezekhez társulnak a műveleti tervezés specifikus paraméterei: parancsnoki szándék és kockázatvállalás, rendelkezésre álló erők (force capabilities), idő- és térbeli korlátok, logisztikai realitások, valamint a jogi/ROE és politikai megkötések. [33], [34], [10], [42]

A gyakorlatban célszerű a MI számára egy explicit „tervezési szerződést” (planning contract) létrehozni: melyek azok a változók, amelyeket a rendszer szabadon variálhat (pl. erők elosztása, ütemezés, LOE-struktúra), és melyek nem változtathatók (pl. tiltott célkategóriák, koalíciós korlátozások, minimum erőszint, időzárak). A tervezési szerződés része a minőségbiztosítás is: minden MI-által javasolt COA esetén kötelezően előállítandó a hivatkozások, feltételezések, bizonytalanságok és várható következmények listája, hogy a törzs azonnal ellenőrizni tudja a javaslat logikáját. Ez a megközelítés egyben csökkenti a „hallucináció” és a rejtett feltételezések kockázatát is. [12], [44]

Az alábbi táblázat a COPD lépések és input/outputk összefüggéseit mutatja.

7-1. táblázat: MI-támogatási funkciók a COPD 3–4. fázis tipikus lépései szerint (saját szerkesztés [33], [34], [42], [12] alapján).

COPD lépés / termék	MI-támogatási cél	Fő inputok	Kimenet / artefaktum	Kockázat és kontroll
COA/MRO vázlatképzés	Gyors alternatíva-generálás, „ötlettér” feltárása	Mission analysis, JIPOE/PMESI I, parancsnoki iránymutatás	COA-vázlatok, LOE/effektus térkép, hipotézisek	Human-in-the-loop; tiltások, kötelező forrásjelölés
COA kidolgozás	Konzisztencia-ellenőrzés, hiányok feltárása	Erőlista, logisztika, idővonal, ROE/politikai korlátok	COA részletezés (fázisok, döntési pontok), erőigény	Validációs checklist; magyarázhatóság ; érzékenységvizsgálat

COA elemzés / wargaming	Szcenáriók generálása, gyors szimulációk, ellenség-reakciók becslése	COA leírás, fenyegetésmodell, környezetváltozók	Kimeneti mutatók (idő, kockázat, erőforrás-fogyás), narratív összefoglaló	Model credibility; VV&A; sztochasztikus futtatások
COA összehasonlítás	Többkritériumos rangsorolás és trade-off elemzés	Mutatók, súlyok, kockázati preferenciák	COA comparison matrix, érzékenységi grafikonok	Átlátható súlyozás; audit trail; döntési napló
CONOPS kidolgozás	Vázlatból koncepció: szöveg és vizuális konzisztencia, mellékletek előkészítése	Kiválasztott COA, parancsnoki döntések, koordináció	CONOPS-fejezetek, főbb ábrák, feladat-felelősség mátrix	Forráskezelés; minősítési szabályok; verziókezelés
OPLAN/OPO RD	Dokumentum-összeállítás, mellékletek/függelékek generálása sablonból	CONOPS, logisztika, C2, kommunikáció, ISR, cyber, IO	Tervdokumentum-csomag, konzisztencia-jelentés	„Content firewall”; titokvédelem; szemantikus ellenőrzés
Frissítés és változáskezelés	Gyors újratervezés: mi változik, mi nem, hatáselemzés	Új hírszerzés, új korlátok, erőváltozások	Delta-jelentés, új COA-variánsok, kockázati frissítés	MLOps/Drift-monitoring; jóváhagyási workflow

7.1.2. COA/CONOPS adatmodell: feladatok, hatások, erőforrások és idő-tér összerendelése

A MI-vezérelt COA-generálás kulcsa a megfelelő reprezentáció. Ha a COA csak narratív szöveggént létezik, akkor a rendszer legfeljebb nyelvi műveleteket (összegzés, átírás, fordítás) tud végezni. Ahhoz azonban, hogy alternatívákat tudjon generálni és értékelni, a COA-nak legalább részben strukturált formában kell megjelennie. A tervezési gyakorlatban ez már ma is jelen van: a COA-t gyakran LOE-k (lines of effort) és műveleti fázisok szerint bontják, és a törzs COA-vázlatokkal, idővonalakkal, döntési pontokkal és erő-hozzárendelésekkel dolgozik. [33], [42]

A fejezet MI-szemlélete szerint a COA/CONOPS minimális adatmodellje négy fő komponens köré építhető: (1) célok és hatások (objectives, desired effects), (2) feladatok és tevékenységek (tasks/activities), (3) erőforrások és képességek (forces/capabilities), valamint (4) idő-tér leképezés (phasing, sequencing, geometry). Egy ilyen modell természetes módon ábrázolható gráfként: a csomópontok lehetnek feladatok, szereplők, objektumok és események, az élek pedig függőségeket (előfeltétel, okozat, támogatás), erő-hozzárendelést vagy információáramlást fejeznek ki. A gráfrepresentáció előnye, hogy egyszerre támogatja a logikai ellenőrzést (nincs-e körkörös függőség, hiányzó előfeltétel), és a keresést/optimalizálást (melyik csomópont-sorozat adja a legjobb hatást a legkisebb erőforrással).

A strukturált reprezentáció gyakorlati megvalósításához célszerű egy „COA-séma” (COA schema) vagy ontológia kialakítása, amely rögzíti a kötelező mezőket. Tipikus kötelező elemek: parancsnoki szándék-összefoglaló; műveleti célok; erők és szerepkörök; fázisok és fő események; döntési pontok és kiváltó indikátorok; kritikus függőségek (logisztika, kommunikáció, engedélyek); valamint mérőszámok (measuring effectiveness és measuring performance). A modellnek kezelnie kell a bizonytalanságot is: nem csak egyetlen „igaz” értéket, hanem tartományokat és valószínűségi feltételezéseket (pl. időtartam eloszlása, ellenséges reakció késleltetése) kell tárolnia. Ez teszi lehetővé a későbbi sztochasztikus szimulációt és érzékenységvizsgálatot.

7-2. táblázat: Javasolt minimál COA/CONOPS adatséma (példa) MI-támogatott generáláshoz (saját szerkesztés [33], [42], [12] alapján).

Adatmodell elem	Tartalom (példamezők) / MI-hasznosítás
Célok (Objectives)	Kívánt végállapot, részcélok; kapcsolódás politikai iránymutatáshoz; mérőszámok előkészítése.
Várt hatások (Desired effects)	Operatív/taktikai hatások; PMESII-PT dimenziók szerinti hatástérkép; másodrendű hatások jelölése.
Feladatok (Tasks)	Tevékenységlista LOE/fázis bontásban; előfeltételek; támogatási viszonyok; kritikus út.
Szereplők és felelősség	Own/ally/partner szereplők; felelősségi mátrix (RACI-szerű); C2 kapcsolatok.

Erők és képességek	Erőelemek, képesség-katalógus; készségi és fenntarthatósági paraméterek; korlátozások.
Idővonal és ütemezés	Fázisok, mérföldkövek, döntési pontok; időtartam-tartományok; szinkronizációs követelmények.
Térbeli leképezés	AOI/JOA/TAI fogalmak; releváns területek; mozgás- és elérési korlátok (nem célpontszinten, hanem absztrakcióban).
Korlátok és tiltások	ROE/politikai korlátok; tiltott hatások; minimum erő/kapacitás; időzárak; adatvédelmi/infosec korlátok.
Feltételezések	Explicit tervezési feltételezések; bizonytalansági szintek; ellenőrzési terv (mikor/ki validálja).
Indikátorok és kiváltók	Esemény-indikátor párok döntési pontokhoz; riasztási küszöbök; helyzetképhez kapcsolás.
Kockázat és mitigáció	Fő kockázatok, valószínűség/hatás; mitigációs opciók; reziliencia elemek.
Metaadat és audit	Források, verziók, felelős személy, időbélyeg; MI-modell verzió; magyarázat és naplózás.

7.1.3. COA-generálási megközelítések: szabály, eset, keresés és generatív MI

A COA-generálás MI-támogatása több technikai paradigmára épülhet. A legegyszerűbb megoldás a sablon- és szabályalapú generálás: a törzs által elfogadott koncepció-sablonok (pl. stabilizációs, elrettentési, védelmi műveleti logikák) és az ezekhez rendelt „if-then” szabályok gyorsan képesek konzisztens COA-vázlatokat előállítani. Előnyük a kontrollálhatóság és az átláthatóság; korlátjuk, hogy a szabálybázis karbantartása munkaigényes, és a ritka, új helyzetekben könnyen „kifogynak” a lefedettségéből.

A második megközelítés az esetalapú következtetés (case-based reasoning): a rendszer korábbi, anonimizált és absztrahált műveleti tervekből, CONOPS-okból vagy OPLAN-szerkezetekből keres hasonló mintákat, majd adaptálja azokat a jelen helyzetre. Ennek előnye, hogy a szervezet bevált megoldásait „újrahasznosítja”, és gyors tanulási görbét ad. Korlátja, hogy a történeti példák torzíthatják a jelen helyzet értelmezését (túlzott analógia), és a hasonlóságmérés (melyik múltbeli eset releváns?) nem triviális.

A harmadik paradigma a klasszikus MI-tervezés (automated planning): a feladatot explicit cél–előfeltétel–hatás struktúrában írja le, majd kereső algoritmusokkal (pl. hierarchikus feladatbontás – HTN, részleges rendezésű tervezés) állít elő tervet. A szakirodalomban több korai katonai alkalmazási kísérlet is megjelent, amelyek a műveleti tervezést mint formális tervezési problémát kezelték (pl. interaktív tervező rendszerek és prototípus AI-plannerek). [49]–[51]

A negyedik, napjainkban leggyorsabban terjedő megközelítés a generatív MI (különösen a nagy nyelvi modellek – LLM-ek) tervezési asszisztensként való alkalmazása. Itt a modell jellemzően nem „oldja meg” a tervezési problémát formális értelemben, hanem a törzs számára olvasható COA-narratívát, CONOPS-fejezeteket, összehasonlító érveket és kockázati listákat generál. A használható megoldásoknál azonban kritikus a tudás-visszakeresés (retrieval) és a forrásalapú generálás (RAG): a modellnek a saját „paraméter-memóriája” helyett a hiteles doktrinális és helyzetadatokra kell támaszkodnia, és minden állítását vissza kell vezetnie konkrét forrásra vagy feltételezésre. [12], [44], [45]

A gyakorlatban a legjobb eredményt jellemzően hibrid megoldások adják: a formális tervezési modul (szabály/HTN) biztosítja a konzisztenciát és a korlátok betartását, míg a generatív modul (LLM) a dokumentumalkotásban, összegzésben és a döntési érvelés strukturálásában segít. A hibrid megközelítés illeszkedik a NIST AI RMF szemléletéhez is: a kockázatos képességeket (pl. autonóm döntés) korlátozni kell, miközben az alacsonyabb kockázatú, nagy hatású funkciókat (pl. dokumentumsablonok kitöltése, változáskövetés) kontrolláltan lehet automatizálni. [12]

7-3. táblázat: COA-generálási megközelítések összehasonlítása (saját szerkesztés [42], [12], [49]–[51] alapján).

Megközelítés	Tipikus technika	Erősség	Gyengeség	Javasolt használat (COPD 3–4)
Sablon- és szabályalapú	SOP-k, doktrinális sablonok, if-then szabályok	Átlátható, könnyen auditálható, gyors	Nehezen skálázódik új helyzetekre; karbantartási teher	COA-vázlatok, OPLAN-sablon kitöltés, konzisztencia ellenőrzés
Esetalapú következtetés	Hasonlóságmérés, mintakeresés	Szervezeti tudás újrahasznosítása; gyors adaptáció	Túlzott analógia; relevancia-	Kampánytervek, bevált LOE minták,

	korábbi tervekben		mérés nehéz; adatminőség függő	konceptió- variánsok
Klasszikus MI- tervezés	HTN, részleges rendezésű tervezés, constraint satisfaction	Formális cél–korlát kezelés; „bizonyítható” konzisztencia	Modellezési költség magas; bemenet formalizálást igényel	Erő- és ütemezési változatok generálása; kritikus út elemzés
Generatív MI (LLM+RAG)	Forrásalapú szöveggenerál ás, összegzés, érvelés- szervezés	Gyors dokumentumalkotá s; kommunikáció; alternatíva-érvelés	Hallucináció kockázat; rejtett feltételezése k; adatvédelmi kockázat	CONOPS/OPLA N narratív részei, briefek, döntési érvelés
Hibrid architektúra	Formális tervező + generatív asszisztens + szimuláció	Korlátbetartás + jó használhatóság; skálázható workflow	Integráció és governance komplex; tesztelési igény magas	COA generálás– wargaming– OPLAN összeállítás teljes láncban

7.1.4. COA-elemzés szimulációval: wargaming, agent-based és sztochasztikus értékelés

A COPD 3. fázisában a COA-elemzés célja, hogy a törzs ne csak a COA „belső logikáját”, hanem a várható következményeket is feltárja: hogyan reagálhat az ellenség; hol vannak a töréspontok; milyen erőforrások válnak kritikus szűk keresztmetszetté; és mely feltételezések a legkockázatosabbak. A háborús játék (wargaming) hagyományosan ezt a feladatot szolgálja – a résztvevők szerepjátékos, szabályokkal támogatott módon járják végig a COA lépéseit, majd rögzítik a megfigyeléseket és a tanulságokat. A modern gyakorlatban a wargaming egyre inkább összekapcsolódik a modellezés–szimuláció (M&S) eszközeivel, különösen akkor, ha sok variánst vagy nagy bizonytalanságot kell kezelni. [33], [42], [52]

Az MI két irányban tudja kiegészíteni a wargaminget. Egyrészt előkészítheti a játékot: automatikusan generálhat „játékszabály-szerű” eseménykártyákat, ellenség-reakció opciókat, és azonosíthatja a kritikus döntési pontokat (branch point) a COA gráfrepresentációja alapján. Másrészt feldolgozhatja a játék eredményeit: a megjegyzésekből és eseménynaplókból strukturált tanulságlistát, kockázati regisztert és COA-módosítási javaslatot állíthat elő, miközben minden megállapítást visszaköt a wargame eseményéhez (traceability). Ez

különösen fontos a döntési brief készítésekor, ahol a parancsnok számára világosan meg kell mutatni, milyen bizonyítékok és milyen bizonytalanságok támasztják alá az ajánlást. [37], [12]

A szimulációs támogatás egyik leghasznosabb formája az agent-based megközelítés, amely a komplex adaptív rendszereket (pl. több szereplő, interakciók, lokális szabályok) képes kezelni. Egy, a komplex műveleti környezetre készült amerikai hadsereg-beli kutatási jelentés például kifejezetten az agent-based szimuláció alkalmazhatóságát vizsgálja konzekvencia-elemzésre, és hangsúlyozza, hogy az ilyen modellek akkor adnak értéket, ha a felhasználó tisztában van a modell feltételezéseivel és a validáció korlátaival. [53]

A szimulációs eredmények megbízhatóságát két dolog határozza meg: (1) a modell hitelessége (model credibility), és (2) a bizonytalanság kezelése. A hitelesség nem csak technikai kérdés; részben eljárási is: világos dokumentáció, ellenőrzés–validálás (VV&A), paraméternyilvántartás és a döntési felhasználás korlátainak rögzítése. A bizonytalanság kezelése pedig azt jelenti, hogy nem egyetlen futtatásra építünk, hanem több futtatásra, több paraméter-szenárióval: például Monte Carlo jellegű értékeléssel megadható, hogy a COA várható teljesítménye milyen szórással változik különböző körülmények mellett. A törzs számára ez azért értékes, mert így láthatóvá válik, melyik COA „robosztus” (sok körülmény mellett is elfogadható), és melyik „éles” (csak szűk feltételek mellett működik).

A COA-értékelés során célszerű a szimulációs kimeneteket több szintre bontani: (a) operatív kimenetek (idő, erőfogyás, logisztikai terhelés), (b) kockázati kimenetek (kritikus esemény bekövetkezési valószínűsége, tartalékok kimerülése), és (c) stratégiai/PMESII hatások (legitimáció, információs tér, infrastruktúra terhelés) – utóbbiaknál természetesen csak absztrakt mutatókkal dolgozva. A többdimenziós kimeneteket a MI képes egységes döntéstámogató felületen összerendezni, és a parancsnoki preferenciák szerint érzékenységvizsgálatot készíteni. [12], [42]

7.1.5. Kimenetek: COA-brief, összehasonlító mátrix és „tervezési emlékezet”

A MI-támogatott COA-folyamat akkor hasznos, ha a kimenetei közvetlenül illeszkednek a törzs döntés-előkészítő termékeihez. Ide tartozik a COA-vázlat és a COA-narratíva; a COA-összehasonlító mátrix; a kockázati regiszter; az erőforrás-igény és fenntarthatósági összegzés; valamint a döntési pontok és kiváltó indikátorok listája. A MI egyik hozzáadott értéke a „tervezési emlékezet” kialakítása: a döntések mögötti érvek, a wargaming során azonosított

problémák, és a módosítások indokai strukturáltan visszakereshetők, így a 4. fázisban az OPLAN összeállítása gyorsabb és kevésbé hibázik. [37], [42]

A tervezési emlékezetnek van egy szervezeti tanulási funkciója is. Ha a COA-k és a wargaming megállapításai egységes adatsémában kerülnek rögzítésre (lásd 7-2. táblázat), akkor később – akár más műveletekben, akár gyakorlatokon – esetalapú következtetésre és doktrinális visszacsatolásra is alkalmas. Ez a megközelítés az összhaderőnemi törzsmunkában is releváns, ahol a tervkészítés minősége nagyban múlik a standardizált eljárásokon és a dokumentált döntéseken. [42], [54]

7.1.6. Átmenet a COA-tól a CONOPS-ig: MI a koncepció kidolgozásában

A COPD 4. fázisának első nagy lépése a kiválasztott COA „átfordítása” koncepcióvá (CONOPS). A CONOPS feladata, hogy a COA logikáját – a célok, a műveleti gondolatmenet és a fő erő kifejtés – olyan formában rögzítse, amely egyszerre érthető a parancsnok számára és végrehajtható a beosztott szinteknek. A JP 5-0 és a NATO tervezési elvek szerint a CONOPS nem egyszerűen részletezés, hanem koherencia-teremtés: a COA „gráf” (feladatok és függőségek) és a tervdokumentum (szöveg, ábrák, mellékletek) közötti megfeleltetés. [33], [34], [42]

MI-támogatásban a CONOPS-kidolgozás egyik legnagyobb értéke a konzisztencia biztosítása. A hibrid tervezési architektúra (7.4.3) esetén a COA strukturált adatsémája már tartalmazza a célokat, LOE-ket, fázisokat, döntési pontokat és erő-hozzárendeléseket. Egy generatív modul (LLM) ebből képes a CONOPS narratív fejezeteit előállítani, de a minőség kulcsa az, hogy a szöveg minden bekezdése visszaköthető legyen a strukturált elemekhez: ha a narratíva új feladatot említ, annak meg kell jelennie a feladatlistában; ha új kockázatot nevez meg, annak szerepelnie kell a kockázati regiszterben. Ezt a megfeleltetést a MI automatizált ellenőrzéssel (sémavalidáció + szemantikus ellenőrzés) tudja támogatni. [12], [42]

A CONOPS tipikus tartalmi elemei – amelyek MI-vel részben automatizálhatók – a következők: (1) a művelet célja és végállapota; (2) a fő erő kifejtés és a műveleti elgondolás; (3) fázisok és átmenetek; (4) C2 és koordináció; (5) kulcsfontosságú képességek (ISR, logisztika, kommunikáció, információs dimenzió, cyber); (6) kockázatok és tartalékok; (7) mérőszámok és jelentési rend (assessment). A generatív eszközök itt különösen a strukturált vázból történő „doktrinális nyelvre” fordításban erősek: röviden, érthetően képesek a törzs

munkáját dokumentumformába önteni, miközben a sablonok biztosítják az egységes szerkezetet. [42], [12]

A MI további hozzáadott értéke a szinkronizációs termékek előállítása. Ilyenek a COA-vázlatból levezetett szinkronizációs mátrixok (melyik LOE melyik fázisban milyen támogatást igényel), a döntési pontokhoz rendelt információigények listája (PIR/FFIR jellegű), valamint a koordinációs feladatok (civil, koalíciós, tárcaközi) áttekintése. Ezek a termékek alapvetően táblázatos/strukturált jellegűek, ezért jól illeszkednek a gépi feldolgozáshoz: a rendszer képes a hiányzó mezők azonosítására, a redundanciák csökkentésére és a cross-reference hibák jelzésére. [37], [42]

7.1.7. OPLAN dokumentum-összeállítás és mellékletek: sablonkitöltés, konzisztencia és biztonság

A COPD 4. fázis végterméke tipikusan az OPLAN/OPORD (hadműveleti terv/parancs) és annak mellékletei (annex) és függelékei (appendix). A tervdokumentumok egyik gyakorlati kihívása, hogy nagy terjedelműek és erősen összefüggnek: egy erő-hozzárendelés változása hatással van a logisztikai számításokra, a kommunikációs tervre, az ISR-prioritásokra és a kockázati mellékletekre is. A manuális szerkesztés ilyen környezetben könnyen inkonzisztenssé válik, különösen időnyomás alatt. A JP 5-0 és a NATO COPD-szemlélet ezért is támaszkodik erősen standardizált szerkezetre és törzseljárásokra. [33], [42], [54]

MI-támogatásban az OPLAN összeállításának „alacsony kockázatú, nagy hasznú” területe a dokumentumsablonok kitöltése és a konzisztencia-ellenőrzés. A rendszer képes lehet a CONOPS strukturált elemeiből automatikusan előállítani a releváns fejezetrészek vázát, majd azokat a törzs szakértői által jóváhagyott sablonok szerint kiegészíteni. Ezzel párhuzamosan futtatható egy szemantikus ellenőrzés: például ugyanaz a feladat ne szerepeljen eltérő megnevezéssel több mellékletben; a fázisok számozása és időrendje egyezzen; és a C2 kapcsolatok konzisztensen jelenjenek meg a műveleti és kommunikációs mellékletben. A módszer lényegi előnye az auditálhatóság: a rendszer jelzi, hogy melyik bekezdés melyik strukturált mezőből származik (trace). [12], [42]

A dokumentum-automatizálás azonban csak akkor felelős, ha a minősítési és információbiztonsági követelmények be vannak építve. A generatív MI-k egyik fő kockázata a nem szándékolt adatkiáramlás (pl. külső felhőszolgáltatóba küldött tartalom) és a későbbi visszakereshetőség hiánya. Ezért a katonai alkalmazásoknál a „content firewall” elvét kell

érvényesíteni: a modell csak olyan környezetben futhat, ahol a hozzáférés és az adatmozgás policy-alapú, a promptok és kimenetek naplózottak, és a minősített tartalom elkülönül. Ez összhangban van a NIST AI RMF kormányzási és biztonsági fókuszával is. [12]

7-9. táblázat: Példák az OPLAN tipikus mellékleteire és a lehetséges MI-támogatásra (saját szerkesztés [42], [33], [55] alapján).

OPLAN-rész/melléklet (példa)	Fő tartalom	Lehetséges MI-támogatás (példa)
Műveleti melléklet (Operations)	Feladatok, fázisok, koordináció, szinkronizáció	Sablonkitöltés COA-sémából; kereszt-hivatkozás ellenőrzés; változáskövetés
Hírszerzési/ISR melléklet (Intelligence/ISR)	PIR/FFIR, felderítési prioritások, gyűjtési terv	Indikátorok hozzárendelése döntési pontokhoz; VOI-vezérelt ISR kiosztási javaslat
Logisztikai melléklet (Logistics)	Ellátási koncepció, készletek, utánpótlás, ütemezés	Fenntarthatósági számítás; robusztus tartalék tervezés; inkonzisztenciák jelzése
Kommunikációs/CIS melléklet	C2 kapcsolatok, kommunikációs terv, hálózati igények	Séma- és terminológia-konzisztencia; hálózati terhelési kockázatok jelzése
Cyber / információs dimenzió	Cyber hatások, védelmi intézkedések, IO narratíva	Jogi/politikai korlát check; időablak-figyelés; integráció a döntési pontokkal
Kockázat és assessment	Kockázati regiszter, mitigáció, mérőszámok	Automatikus kockázati lista összevezetés wargame naplóból; dashboard; bizonytalanság jelölés

7.2. Optimalizáló algoritmusok erőforrás-elosztáshoz

A COPD 3–4. fázisában az erőforrás-elosztás nem pusztán „logisztikai” kérdés, hanem a COA életképességének egyik fő determinánsa. A COA akkor megvalósítható, ha a feladatokhoz rendelt erők és képességek időben és térben rendelkezésre állnak, fenntarthatók, és a kritikus támogatások (pl. ISR, kommunikáció, logisztika, egészségügyi biztosítás, cyber/infó műveletek) összehangoltan működnek. A 4. fázisban a terv mellékletei és függelékei konkrét erő- és eszköz-hozzárendelést, ütemezést és fenntarthatósági számításokat igényelnek, ezért az optimalizáció természetes kapcsolódási pont a CONOPS–OPLAN átmenetben. [33], [42]

7.2.1. Optimalizációs problémátípusok a tervezésben

A katonai tervezésben az optimalizáció tipikusan több, egymással összefüggő problémátípusban jelenik meg. Leggyakoribbak: (a) hozzárendelési problémák (assignment) – mely erőelem mely feladatot támogatja; (b) ütemezési problémák (scheduling) – mikor melyik képesség áll rendelkezésre, hogyan férnek el a fázisokban; (c) útvonal- és ellátási problémák (routing, distribution) – hogyan jut el az utánpótlás a fogyasztási pontokra; (d) készlet- és fenntarthatósági problémák (inventory, sustainment) – milyen ütemben fogy a készlet és hol a töréspont; (e) hálózati problémák (network flow) – kommunikációs/logisztikai hálózatok terhelése és sérülékenysége. E problémák közös jellemzője, hogy sok korláttal (kapacitás, idő, kompatibilitás, kockázat) és gyakran több célfüggvénnyel (hatásosság vs. veszteség vs. költség) dolgoznak.

7.2.2. Determinisztikus, robusztus és sztochasztikus optimalizáció

A hagyományos (determinisztikus) optimalizáció – például lineáris programozás (LP) vagy egészértékű vegyes programozás (MILP) – akkor használható jól, ha a bemenetek (erő, idő, kapacitás) kellően pontosan becsülhetők, és a korlátok egyértelműen formalizálhatók. A katonai tervezésben azonban a bizonytalanság nem kivétel, hanem norma: ezért a robusztus és a sztochasztikus optimalizáció különösen releváns. Robusztus megközelítésnél a megoldást úgy keressük, hogy az a bemeneti tartományon belül „elég jó” maradjon; sztochasztikus megközelítésnél pedig a bemeneteket eloszlásokkal írjuk le, és a várható érték mellett a kockázati mérőszámokat (pl. worst-case, CVaR, küszöb alatti teljesítmény valószínűsége) is figyelembe vesszük. Tervezési szempontból ez közvetlenül támogatja a parancsnoki kockázatvállalás explicitté tételét: a COA nem csak „optimális”, hanem „vállalható” is lesz.

7.2.3. Heurisztikák, metaheurisztikák és megerősítéses tanulás

Nem minden tervezési probléma formalizálható jól LP/MILP formában, illetve sok probléma esetén az egzakt megoldás túl lassú. Ilyenkor heurisztikák és metaheurisztikák (pl. genetikusan algoritmus, szimulált hűtés, tabu-keresés) adhatnak gyors, közelítő megoldást. A metaheurisztikák előnye, hogy rugalmasan kezelik a nemlineáris célfüggvényeket és a komplex korlátokat is; hátrányuk, hogy a megoldás minősége és reprodukálhatósága erősen

függ a paramétereiktől. A megerősítéses tanulás (reinforcement learning – RL) és az approximate dynamic programming a dinamikus erőforrás-elosztásnál (pl. adaptív ISR-kiosztás, változó logisztikai igények) lehet ígéretes, de katonai felhasználásban különösen fontos a biztonságos tanulás és a szabályozott környezet (szimuláció) – éles alkalmazásban az RL „felfedezési” fázisa nem megengedhető. [12], [44]

7.2.4. Optimalizáció és MI-kormányzás: átláthatóság, audit és felelős használat

Az optimalizáló algoritmusok eredményei – különösen ha több célfüggvény és bizonytalanság is jelen van – könnyen „fekete dobozként” jelenhetnek meg a döntéshozó előtt. Ez a katonai tervezésben kockázatos: a parancsnoknak értenie kell, miért javasol a rendszer egy adott erő-elosztást, mely korlátok aktívak, és milyen trade-offok történnek. A NIST AI RMF megközelítése alapján ezért az optimalizációs modulokat is úgy kell kialakítani, hogy rendelkezzenek: (a) dokumentált célfüggvénnyel és súlyokkal, (b) érzékenységvizsgálattal (mi változik, ha egy feltétel romlik), (c) magyarázható kimenettel (kritikus korlátok, shadow price jellegű jelzések), és (d) auditnaplóval (bemenetek, verziók). A NATO felelős MI irányai szintén azt erősítik, hogy a tervezésben az MI csak akkor fogadható el, ha a döntési felelősség és az elszámoltathatóság megmarad. [12], [45]

7-4. táblázat: Optimalizációs módszerek áttekintése és tipikus tervezési alkalmazások (saját szerkesztés [42], [12], [55] alapján).

Módszer	Tipikus tervezési feladat	Előny	Korlát	Gyakorlati kimenet (példa)
LP/MILP (egzakt)	Erők hozzárendelése, ütemezés, készlet-tervezés	Optimális megoldás (modell szerint); erős kontroll a korlátokon	Modellezési igény magas; bizonytalanság kezelése korlátozott	Erő-hozzárendelési táblázat, ütemterv, fenntarthatósági számítás
Robusztus / sztochasztikus	Tartalék-tervezés, ellátási bizonytalanság kezelése	Reziliensebb terv; kockázat explicit kezelése	Számítási igénye magas; eloszlásbecslés szükséges	Kockázati sávok; „worst-case” és várható érték összehasonlítás

Metaheurisztikák	Komplex ütemezés és routing, nemlineáris célok	Gyors, rugalmas; nagy keresési tér	Nincs garancia optimálisra; paraméter-függő	Közelítő ütemterv; alternatív megoldások halmaza
RL / dinamikus döntés	Adaptív erőforrás-kiosztás változó helyzetben	Tanulhat környezetből; képes nemlineáris politikára	Biztonságos tanulás nehéz; validáció kritikus	Döntési szabály/politika; riasztási küszöbök; adaptív kiosztási javaslat
Hibrid (opt.+szimul.)	COA-k összehasonlítása bizonytalanság mellett	Optimalizálás + következményvizsgálat egyben	Integráció bonyolult; adatigény nagy	Trade-off görbék, Pareto-front; COA ranking + kockázat

7.2.5. Többcélú optimalizáció és Pareto-szemlélet a COA-választásban

A valós tervezés ritkán egycélú. Tipikusan egyszerre kell maximalizálni a hatásosságot (pl. célok elérése, idő), miközben minimalizálni kell a veszteséget, a költséget és a politikai kockázatot. A többcélú optimalizáció egyik hasznos fogalma a Pareto-határ: azok a megoldások, amelyeknél egy cél javítása csak egy másik cél romlásával érhető el. MI-támogatásnál a cél nem feltétlenül egyetlen „legjobb” megoldás kiszámítása, hanem a Pareto-front bemutatása a parancsnok számára: milyen opciók érhetők el különböző kockázatvállalási szinteken. Ez a megközelítés illeszkedik a COA-összehasonlítás céljához, ahol a törzsnek nem csak rangsorolnia kell, hanem a döntés érveit is meg kell fogalmaznia. [42]

7.2.6. Erőforrás-elosztás különleges területei: ISR, információs műveletek és cyber

Kiemelt terület az ISR (intelligence, surveillance, reconnaissance) képességek kiosztása. Az ISR-eszközök (szenzorok, felderítő platformok, elemző kapacitások) korlátozottak, miközben a COA-elemzés és a végrehajtás során is versenyeznek értük a feladatok. Optimalizációs szempontból az ISR-kiosztás gyakran többperiódusú hozzárendelési probléma, amelyben a „haszon” nem csak információmennyiség, hanem információérték (value of information): mennyivel csökkenti a bizonytalanságot egy döntési pont előtt. MI-eszközökkel a VOI becslése

összekapcsolható prediktív modellekkel és a fenyegetésindikátorokkal, így az ISR nem „mindent mindenkor” elv alapján, hanem célzottan támogathatja a COA-t. [10], [42]

Hasonlóan összetett a cyber és az információs dimenzió erőforrás-kezelése. A cyber műveletek tervezési útmutatói hangsúlyozzák a cél–eszköz–hatás összhangját, valamint a jogi és politikai korlátok meghatározó szerepét. A cyber képességek gyakran „időérzékenyek” (window of opportunity), és a hatásuk nehezen mérhető, ezért a tervezésben különösen fontos az indikátorokhoz kötött döntési pontok és a tartalék opciók kidolgozása. Optimalizációs szemszögből ez tipikusan robusztus, többcélú problémaként kezelhető, ahol a cél a kockázat kontrollálása, nem pedig a hatás maximalizálása. [55], [12]

7.2.7. Érzékenységvizsgálat, „what-if” elemzés és red teaming

Az optimalizációs eredmények tervezési értéke nagyban függ attól, hogy a törzs képes-e megérteni a megoldás „törékenységet”. Ezért a 3–4. fázisban az optimalizációt érdemes rendszeresen érzékenységvizsgálattal kiegészíteni: mely korlátok a legszorosabbak (bottleneck), mely paraméterek változására ugrik a megoldás (pl. készletcsökkenés, időzár, szállítási kapacitás), és milyen tartalék opciók csökkentik a kockázatot. A MI különösen erős a „what-if” kérdések gyors futtatásában és strukturált bemutatásában: például egy parancsnoki döntési pontnál a rendszer előállíthatja, hogy a legvalószínűbb és a legveszélyesebb szcenárióban hogyan változik az erőforrás-elosztás, és melyik LOE kerül veszélybe. [42], [12]

Az érzékenységvizsgálat red teaming szempontból is hasznos: a red cell kérdései gyakran arra irányulnak, hogy „mi történik, ha az ellenség ezt vagy azt teszi”, vagy ha egy kritikus infrastruktúra kiesik. Ezek a kérdések formálisan a korlátok és paraméterek módosításaként jelennek meg. Ha az optimalizációs/szimulációs hurok jól integrált, akkor a red teaming nem ad hoc vita, hanem gyors, dokumentált elemzés lehet – miközben továbbra is érvényes, hogy a modell csak a modell; ezért a kimenetek bizonytalanságát és korlátait mindig kommunikálni kell. [37], [12]

7.2.8. Adatminőség és interoperabilitás: az optimalizáció „üzemeltethetősége”

Az erőforrás-elosztó modulok gyakorlati bevezetésének legnagyobb akadálya ritkán a matematikai algoritmus, sokkal inkább az adatminőség és az interoperabilitás. A tervezés során az erő-adatok, készletek, készenléti állapotok, szállítási kapacitások és korlátozások több

rendszerben, eltérő formátumban léteznek. Ha nincs egységes fogalmi modell és adatkapu (data governance), az optimalizáció „szép” eredményt adhat, de nem lesz végrehajtható, mert a bemenetek nem tükrözik a valós képességállapotot. Ezért a MI-alapú optimalizációt célszerű úgy kezelni, mint egy üzemszerűen működtetett döntéstámogató szolgáltatást: bemeneti adatvalidációval, verziózott adatforrásokkal, és a kimenetek jóváhagyási folyamatával. Ez közvetlenül kapcsolódik a NIST AI RMF „Manage” funkciójához (folyamatos monitorozás és drift-kezelés). [12]

7.3. Integráció háborús játékok eszközeivel

A háborús játék (wargaming) a COA-elemzés egyik központi eszköze: egyszerre teszteli a COA belső koherenciáját, az ellenség várható reakcióit és a törzs gondolkodási fegyelmét. A COPD és a JP 5-0 logikája szerint a wargaming nem csak „szimuláció”, hanem strukturált beszélgetés és döntési tanulás is: a résztvevők a COA lépéseit végigjátszva azonosítják a kritikus feltételezéseket, a kockázatokat és a szükséges döntési pontokat. [33], [42]

7.3.1. MI a wargaming előkészítésében: szcenárió, eseménykártya, ellenség-modellezés

A wargaming előkészítése jelentős törzsidőt igényel: szabályrendszer kiválasztása, szerepek és „red cell” felállítása, kezdeti helyzet meghatározása, valamint a COA és az ellenség lehetséges reakcióinak áttekintése. MI-eszközök itt kétfajta támogatást adhatnak. Az első a szcenárió-generálás: a rendszer a JIPOE-termékek és a fenyegetésvértékelés alapján több, egymástól relevánsan eltérő forgatókönyvet (most likely / most dangerous jellegű variánsokat) állíthat elő. A második az ellenség-viselkedés támogatása: nem „automatikus ellenséggként”, hanem a red cell munkáját segítő opciókészletként, amely rávilágít arra, milyen reakciók következnek logikusan az ellenség célrendszeréből. [10], [42], [37]

7.3.2. MI a wargaming levezetésében: félautomata adjudikáció és eseménynapló

A wargaming során az adjudikáció (a lépések következményeinek megállapítása) gyakran vitás pont: a résztvevők eltérően ítélik meg, mi „reális”. Félautomata megoldásoknál a MI nem a végső döntőbíró, hanem egy konzisztencia-ellenőrző és „emlékeztető” rendszer: jelzi, ha egy lépés ellentmond a korábban rögzített korlátoknak, vagy ha egy következmény valószínűsége

a megadott paraméterek szerint alacsony. A cél a transzparens deliberáció támogatása, nem pedig a vita lezárása. A wargaming során keletkező eseménynapló (lépés, reakció, eredmény, megjegyzés) különösen értékes adatszempontból; MI-vel automatizálható a napló strukturálása, címkézése és a tanulságok kinyerése. A wargame validációjával foglalkozó tanulmányok rámutatnak, hogy a módszer értéke erősen függ attól, milyen fegyelmezetten rögzítik és elemzik az eredményeket. [52], [37]

7.3.3. M&S eszközök és MI: adatsere, digitális iker és kimeneti analitika

A modern M&S eszközök (agent-based, diszkrét esemény-szimuláció, rendszer-dinamika) képesek a wargaminget kvantitatív kimenetekkel támogatni, de a kapcsolódás sokszor adatintegrációs problémába ütközik. Ha a COA struktúrája, a környezet paraméterei és a mérőszámok nem egységesek, akkor a modellfuttatások eredményei nehezen hasonlíthatók össze. A 7-2. táblázatban bemutatott minimál adatséma ezért itt is kulcs: a COA-t és a környezetet ugyanabban a fogalmi térben kell leírni, amelyet a szimuláció „megért”. Az agent-based megközelítésről szóló jelentések kiemelik, hogy a paraméterek és szabályok transzparens dokumentációja nélkül a modell félrevezető biztonságérzetet adhat. [53]

Technikai szempontból célszerű egy „digitális iker” (digital twin) jellegű adatmagot létrehozni a tervezéshez: nem a valóság teljes leképezését, hanem a tervezéshez releváns entitásokat, kapcsolatrendszereket és állapotváltozókat (pl. erők, infrastruktúra állapota, kommunikációs csatornák elérhetősége, indikátorok) egységes adatmodellben. Ebbe a magba csatlakozhat a szimulátor (input/paraméter), a wargaming napló (eseményadat), és a MI-analitika (értékelés, anomália, trend). Az integráció fő kihívása a jogosultság és minősítés kezelése: a tervezési adatmag egyszerre tartalmazhat nyílt és minősített elemeket, ezért a rendszerarchitektúrának „tartalom-tűzfalat” kell alkalmaznia (policy-alapú hozzáférés). [12]

7.3.4. Szervezeti és emberi tényezők: bizalom, torzítás és standardok

A wargaming és a szimuláció MI-támogatása csak akkor fog beépülni a törzs rutinjába, ha a felhasználók bíznak az eszközben, de nem vakon. A túlzott bizalom (automation bias) és a túlzott szkepszis (automation disuse) egyaránt káros: az előbbi elfedi a hibákat, az utóbbi elpazarolja a képességnövelő potenciált. Ezért fontosak az analitikai standardok és a minőségbiztosítási lépések. Az ICD 203 analitikai standardjai – bár hírszerzési termékekre

készültek – jó analógia a tervezési analitikához is: következetesség, forrásjelölés, alternatívák megvizsgálása, bizonytalanság kommunikálása. Ezeket a szabályokat a MI-vel támogatott wargaming kimeneteire is alkalmazni kell. [37]

7-5. táblázat: MI-integrációs minták wargaming és M&S eszközökhöz (saját szerkesztés [33], [42], [53], [52], [12] alapján).

Integrációs pont	Adat (bemenet/kimenet)	MI-funkció	Előny	Korlát / kontroll
Előkészítés (pre-game)	JIPOE, COA-séma, fenyegetésmodell	Szenárió- és eseménykártya generálás; kritikus pontok azonosítása	Gyorsabb felkészítés; variánsok biztosítása	Forrásalapúság; red cell validáció
Levezetés (in-game)	Lépés–reakció–eredmény eseménynapló	Konzisztencia-ellenőrzés; félautomata adjudikáció-javaslat	Kevesebb „ad hoc” döntés; átlátható vita	Emberi döntőbíró; modellkorlátok kommunikálása
Útóelemzés (post-game)	Napló + megjegyzések + mutatók	Tematikus kinyerés; kockázati regiszter; COA-módosítás javaslat	Gyors tanulásgképzés; jobb dokumentáltság	ICD 203-szerű standard; audit trail
Szimulációs hurok	COA struktúra + paraméterek + futtatások	Batch futtatás menedzsment; Monte Carlo; érzékenység	Robusztusság látható; kockázati sávok	VV&A; paraméter-fegyelem; drift kontroll
Tervdokumentum hurok	Tanulások + módosítások + verziók	CONOPS/OPLA N frissítés, változáskövetés, sablonkitöltés	Kevesebb adminisztratív hiba; gyors frissítés	Minősítés/infoc; verziókezelési jóváhagyás

7.3.5. Információbiztonság és adatkezelés a wargaming–MI integrációban

A wargaming és M&S eszközök MI-támogatott integrációja különös információbiztonsági kihívásokat vet fel, mert többféle adat keveredhet: minősített helyzetinformációk, saját erő képességei, szövetségi megkötések, valamint nyílt forrású (OSINT) elemek. Ha a MI

komponens nem megfelelően van izolálva, fennáll a kockázat, hogy minősített tartalom kerül nem megfelelő környezetbe, vagy hogy a kimenetekben nem szándékoltan felfedődik érzékeny információ. E kockázat kezelésére a „szeparált adatcsatornák” és a policy-alapú hozzáférés elve alkalmazható: a MI csak olyan adatot érhet el, amelyhez a felhasználó is jogosult, és a minősítési szintet a kimenetekben automatikusan jelölni kell. A naplózás és audit itt nem adminisztratív teher, hanem a felelős használat feltétele. [12]

Gyakorlati mitigáció lehet a szintetikus (synthetic) adatok használata a fejlesztés és tesztelés során. A szintetikus adatok lehetővé teszik, hogy a wargaming–MI integráció technikai komponenseit (adatsere, logolás, elemzés) éles környezethez hasonló adatszerkezettel, de érzékeny tartalom nélkül lehessen kipróbálni. A **fejezet 7.4.4. demonstrációs** szimulációja is ezt a logikát követi: a paraméterek absztraktak, a módszer azonban alkalmas arra, hogy később, megfelelő jogosultsági és minősítési keretek között valós adatokkal is működjön. [12]

7.4. Empirikus elemzés: prototípusok és szimulációk

A fejezet empirikus része két célt szolgál. Egyrészt bemutatja, hogy a COA–CONOPS–OPLAN lánc MI-támogatása milyen prototípus-architektúrában valósítható meg (adatmodell, komponensek, workflow). Másrészt egy egyszerű, reprodukálható szimulációs példán keresztül illusztrálja, hogyan alakítható a COA-elemzés többfuttatásos (sztochasztikus) értékeléssé, amely számszerűsíti a bizonytalanság hatását. A bemutatott számok nem valós műveleti adatok, hanem demonstrációs paraméterek; a módszer célja a megközelítés szemléltetése, nem konkrét műveleti következtetések levonása.

7.4.1. Kutatási megközelítés és prototípus-logika

A COPD 3–4. fázisát támogató MI-eszközök fejlesztésére a „design science” jellegű megközelítés alkalmas: a kutatás egy konkrét problémakört (COA generálás, erőforrás-optimalizáció, wargaming integráció) egy prototípus artefaktum létrehozásán keresztül vizsgál, majd az artefaktumot mérhető kritériumok mentén értékeli. A prototípusfejlesztés előnye, hogy a doktrinális elvárások és a technikai lehetőségek közti kompromisszumok korán felszínre kerülnek: mely adatok hiányoznak, hol törik meg az automatizáció, és milyen kontrollok szükségesek a felelős használatához. [12], [42]

A prototípus értékeléséhez a fejezet a következő minőségi dimenziókat használja: (a) hasznosság (a törzs munkáját ténylegesen gyorsítja-e), (b) konzisztencia és korláttartás, (c) magyarázhatóság és auditálhatóság (miért ezt javasolja), (d) robusztusság bizonytalanság mellett, (e) információbiztonság és adatvédelem, valamint (f) integrálhatóság a meglévő tervezési és M&S eszközökkel. Ezek a dimenziók összhangban vannak a NIST AI RMF által javasolt kockázati gondolkodással. [12]

7.4.2. Prototípus-architektúra: COA-generátor, optimalizáló modul és szimulációs hurok

A COA–CONOPS–OPLAN MI-támogatás prototípus-architektúrája három fő technikai rétegre bontható. Az első az adat- és tudásréteg: ide tartozik a COA-séma/ontológia (7-2. táblázat), a doktrinális sablonok, a műveleti környezetből származó strukturált adatok (erők, korlátok, idővonalak), valamint a hivatkozott források tárháza. A második az analitikai réteg: COA-generátor (hibrid: szabály/HTN + LLM), erőforrás-optimalizáló modul, és szimuláció/Monte Carlo motor. A harmadik a felhasználói és kormányzási réteg: tervezői felület, verziókezelés, jogosultságkezelés, naplózás és audit. A rétegek közti fő interfész a strukturált COA-reprezentáció; minden komponens ezt olvassa és írja, így a rendszer modularizálható és tesztelhető. [12], [42]

A prototípus működését egy tipikus munkafolyamat írja le: (1) a törzs rögzíti a parancsnoki iránymutatást és a tervezési szerződést (7.1.1), (2) a COA-generátor több változatot állít elő (LOE/fázis bontás), (3) a validációs modul kiszűri a nyilvánvaló ellentmondásokat (hiányzó előfeltételek, tiltott hatások), (4) az optimalizáló modul „összeállítja” az erőforrás-hozzárendelést és ütemezést, (5) a szimulációs modul többfuttatásos értékelést készít, (6) a rendszer COA-összehasonlító mátrixot és kockázati összegzést generál, (7) a döntés után a kiválasztott COA automatikusan „átfordítható” CONOPS/OPLAN dokumentumszerkezetbe sablonok segítségével. A folyamat minden lépésében kötelező a magyarázat és a forrásjelölés, ellenkező esetben a kimenet nem engedhető a döntési briefbe. [37], [12]

7.4.3. Reprodukálható szimulációs példa: két COA összehasonlítása bizonytalanság mellett

A következő egyszerű példa a 7.1.4. pontban tárgyalt sztochasztikus értékelést szemlélteti. Két, absztrakt módon megfogalmazott COA-t hasonlítunk össze (COA-A és COA-B). A COA-A gyorsabb ütemezésre és magasabb intenzitásra épül, míg a COA-B fokozatosabb, több

koordinációt feltételez. A környezet bizonytalanságát három, demonstrációs jellegű változó írja le: (a) ellenség aktivitási szint, (b) időjárás/terep miatti késedelem, (c) általános „friction” (szervezeti és technikai súrlódás). A kimenetek: műveleti időtartam, absztrakt erőforrás-fogyás és egy 0–1 közötti kockázati score. A szimuláció 20 000 futtatásból áll; a cél nem valós értékek előállítása, hanem a módszer demonstrálása.

7-7. táblázat: Monte Carlo jellegű COA-összehasonlítás (demonstrációs paraméterek, saját számítás).

Mutató	COA-A átlag	COA-A p90	COA-B átlag	COA-B p90
Időtartam (nap)	23.54	25.98	27.02	29.49
Erőforrás-fogyás (absztrakt egység)	302.64	328.74	306.93	331.46
Kockázati score (0–1)	0.44	0.60	0.36	0.48
P(Időtartam $\leq$ 26 nap)	0.90	-	0.31	-
P(Kockázat $\geq$ 0.5)	0.32	-	0.07	-

A 7-7. táblázatból látható, hogy a demonstrációs paraméterezés mellett a COA-A átlagosan rövidebb időtartamot eredményez, ugyanakkor magasabb kockázati score-al jár, míg a COA-B hosszabb, de kockázati szempontból „stabilabb” eloszlást mutat. A táblázatban szereplő valószínűségi mutatók (pl. $P(\text{Időtartam} \leq 26 \text{ nap})$, $P(\text{Kockázat} \geq 0.5)$) különösen hasznosak döntési támogatásra: a parancsnoki preferencia (gyorsaság vs. kockázat) expliciten összeköthető a COA-k teljesítményeloszlásaival. A módszer gyakorlati haszna nem az, hogy „megmondja a választ”, hanem hogy a vita fókuszát a feltételezésekre és a robusztusságra helyezi: melyik változó (ellenség aktivitás, időjárás késedelem, friction) okozza a legnagyobb szórást, és milyen mitigáció csökkenti azt.

MI-támogatásban a sztochasztikus értékelés két módon integrálható. Egyrészt a MI automatizálhatja a futtatások menedzsmentjét (batch), a paraméter-szenáriók generálását és a kimenetek összegzését. Másrészt a MI képes „magyarázó” narratívát készíteni az eloszlásokhoz: nem csak azt közli, hogy a COA-A gyorsabb, hanem azt is, hogy „gyors, de érzékeny” (például a kedvezőtlen időjárás késedelemre), és megjelöli, mely tervezési döntési pontokon érdemes tartalék opciót kialakítani. Ez a fajta magyarázat csak akkor elfogadható, ha

a rendszer a megállapításokat visszavezeti a konkrét futtatási eredményekre (traceability) és a bemeneti eloszlásokra. [12]

7.4.4. Erőforrás-optimalizáció demonstráció: VOI-alapú ISR-kiosztás

A 7.2.6. pontban jeleztük, hogy az ISR erőforrások kiosztása jól modellezhető „value of information” (VOI) szemlélettel: nem az a cél, hogy mindenhol legyen felderítés, hanem az, hogy a döntési pontok előtt a bizonytalanság a lehető legnagyobb mértékben csökkenjen. Az alábbi, szintetikus példa három ISR-eszköz és öt döntési pont viszonyát mutatja be. Minden ISR-eszköz csak bizonyos döntési pontokat tud lefedni az adott időablakban (kompatibilitási korlát), és minden eszköznek van „költsége” (pl. repült idő, elemzői kapacitás, kitettség). A célfüggvény: maximalizálni a lefedett döntési pontok VOI-értékének összegét a költségek levonása után. A példa kis mérete miatt a „legjobb” hozzárendelés teljes körű enumerációval is meghatározható; nagyobb méretben ugyanez LP/MILP vagy heurisztika formájában futtatható. [42]

7-10. táblázat: Szintetikus VOI-alapú ISR-kiosztás eredménye (demonstráció, saját számítás).

Eszköz	Hozzárendelt DP	VOI	Költség	Nettó érték (VOI–költség)
ISR-A	DP3	6	2.00	4.00
ISR-B	DP2	7	1.00	6.00
ISR-C	DP1	9	1.50	7.50
Összesen				17.50

A táblázat azt illusztrálja, hogy a kompatibilitási korlátok miatt nem feltétlenül az a legjobb, ha a legnagyobb VOI-értékű döntési pontot automatikusan a „legnagyobb” eszközhöz rendeljük. A megoldás értéke a tervezésben az átláthatóság: a rendszer nem csak kiosztást ad, hanem megmutatja a nettó hasznot és a választás okát (korlátok, költségek). Így a törzs gyorsan megvitathatja, hogy a VOI-súlyok megfelelőek-e, vagy hogy egy döntési pont értéke változott-e a helyzetkép frissülésével. A módszer ugyan absztrakt, de jól kapcsolható a COPD 3–4. fázisában megjelenő PIR/FFIR és döntési pont logikához. [10], [42]

7.4.5. Prototípusok értékelése: javasolt mérőszámok és kontrollok

A prototípusok értékeléséhez célszerű olyan mérőszámokat választani, amelyek egyszerre technikaiak és szervezeti jellegűek. Technikai mérőszám például: korlátsértések száma (a MI javaslati közül mennyi sért tiltást), időköltés (mennyi idő alatt generál 3–5 COA-t), stabilitás (ugyanarra a bemenetre mennyire reprodukálható), és robusztusság (a COA rangsor mennyire változik, ha a bemenetek romlanak). Szervezeti mérőszám: a törzs által elfogadott javaslatok aránya, a wargaming utáni módosítások száma és minősége, a döntési brief elkészítési ideje, valamint a visszakereshetőség (audit) minősége. A NIST AI RMF alapján mindehhez hozzárendelhetők kockázati kontrollok: hozzáférés, adatminőség, red teaming és folyamatos monitorozás. [12]

7-8. táblázat: *Javasolt prototípus-mérőszámok és kontrollok a COPD 3–4. fázis MI-támogatásához (saját szerkesztés [12], [42], [37] alapján).*

Dimenzió	Mérőszám (példa)	Mérés/adatforrás	Kontroll (példa)
Korláttartás	Tiltott hatások/korlátsértések száma	COA-séma validáció, szabálmotor log	Kényszerkorlátok; kötelező ellenőrzési lista
Hasznosság	COA-k előállítási ideje; brief elkészítési idő	Tervezési workflow napló	Folyamatmérés; „stop/go” kapuk
Magyarázhatóság	Forrásjelöltség aránya; magyarázat lefedettség	RAG hivatkozások; auditnapló	Kötelező hivatkozás; bizonytalanság jelölése
Robusztusság	Rangsor-stabilitás; érzékenységi mutatók	Monte Carlo futtatások, szcenárió-készlet	Minimum futtatásszám; worst-case riport
Biztonság	Minősítési incidensek; adatkiáramlási kockázat	Hozzáférési log, DLP jelzések	Policy-alapú hozzáférés; „content firewall”
Integráció	API-hibaaarány; adatkonverziós veszteség	Eszköz-integrációs tesztek	Séma-verziózás; interfész-szerződés

7.4.6. Összegző következtetések: mit ad hozzá a MI a COPD 3–4. fázisban, és hol vannak a korlátok?

A fejezet alapján a MI legerősebb hozzájárulása a COPD 3–4. fázisában a tervezési tér gyors feltérképezése és a bizonytalanság kvantifikálása. A COA-generálásnál a MI képes több

variánst gyorsan előállítani, miközben a strukturált adatséma és a korlátrendszer biztosítja a kontrollt. Az optimalizációs modulok támogatják a megvalósíthatóság és fenntarthatóság elemzését, valamint a trade-offok explicitté tételét. A wargaming és szimuláció integrációja pedig lehetővé teszi, hogy a COA-elemzés ne csak szakértői intuícióra, hanem többfuttatásos, dokumentált értékelésre támaszkodjon. [33], [42], [12]

Ugyanakkor a korlátok legalább ilyen fontosak. A MI csak annyira jó, amennyire a bemeneti adatok, a fogalmi modell és a kormányzási keretek engedik: ha a COA nincs formalizálva, ha a korlátok nincsenek rögzítve, vagy ha a rendszer nem biztosít forrásalapú magyarázatot, akkor a MI „gyorsan gyárt” ugyan, de nem megbízhatóan. A felelős katonai alkalmazáshoz ezért szükséges: (a) világos tervezési szerződés, (b) emberi ellenőrzési pontok, (c) auditálhatóság, (d) biztonsági kontrollok, és (e) folyamatos validáció. A következő fejezetek számára ez a logika biztosítja az átmenetet a tervből a végrehajtás támogatásába és a tanuló visszacsatolásba.

8. FEJEZET MI TÁMOGATÁS A COPD 5-6 FÁZISÁBAN: VÉGREHAJTÁS, ÁTMENET (EXECUTION, TRANSITION)

Az 5. (Execution) és a 6. (Transition) fázis a NATO Comprehensive Operations Planning Directive (COPD) logikájában a tervezés „valóságpróbája”: az elfogadott koncepció és műveleti terv (CONOPS/OPLAN) végrehajtása, majd a művelet átadása, lezárása vagy következő műveleti szakaszba történő átvezetése. A döntéshozatal fókusza ebben az időszakban a végrehajtás ütemezésén, az erők és képességek tényleges alkalmazásán, a műveleti hatások mérésén, valamint a változó helyzethez való gyors igazodáson van.[56]

AJP-5 hangsúlyozza, hogy a tervezés és a végrehajtás nem élesen elválasztott, hanem folyamatosan összekapcsolt tevékenység: a végrehajtás során beérkező információk (felderítés, jelentések, civil szereplők visszajelzései) módosíthatják a feltételezéseket, új kockázatokat hozhatnak felszínre, és szükségessé tehetik a terv adaptív újratervezését (branches/sequels, contingency planning).[34]

A 2022-es NATO Strategic Concept a nagyfokú bizonytalanságot, a technológiai versenyt, valamint az információs térben zajló műveletek és befolyásolási tevékenységek jelentőségét emeli ki. Ezzel összhangban a NATO digitális transzformációs törekvései a többdoménű műveletek, az interoperabilitás és a data-driven döntéshozatal feltételeinek megteremtését célozzák.[57], [58]

A mesterséges intelligencia (MI) a COPD 5-6 fázisában elsősorban nem autonóm döntéshozó, hanem döntéstámogató infrastruktúra: gyorsítja a jelentésáramlást és az információfeldolgozást, támogatja a tervvalidációt és a végrehajthatóság folyamatos ellenőrzését, segíti a valós idejű monitorozást (helyzetkép, hatásmérés, kockázati indikátorok), valamint automatizálja a művelet utáni elemzés (after-action review és lessons learned) egy részét. A fejezet célja, hogy a fenti képességeket a doktrinális keretekhez illesztve mutassa be, különös tekintettel a felelős alkalmazásra, az interoperabilitásra és az adatbiztonságra.[12], [21]

8.1 MI a tervvalidációhoz és az adaptív újratervezéshez

A végrehajtási fázisban a törzs számára a legkritikusabb kérdés az, hogy az elfogadott terv (OPLAN és az alárendelt mellékletek/függelékek) továbbra is összhangban áll-e a valósággal. A „terv” ebben az értelemben nem statikus dokumentum, hanem egy feltételezéseken,

erőforrásokon, időzítésen és döntési pontokon alapuló rendszer. A végrehajtás során a helyzet dinamikusán változik: az ellenség alkalmazkodik, a civil környezet reagál, a logisztikai lánc sérülhet, és új politikai korlátozások jelenhetnek meg. Ezért a tervvalidáció a COPD 5. fázisában folyamatos tevékenységként értelmezendő, amely szoros kapcsolatban áll az operatív értékeléssel (assessment) és az ISR-eredményekkel.[34], [59]

A klasszikus validációs módszerek (törzsgyakorlat, harcvezetési ritmus, ellenőrző listák, kézi modellalkotás) idő- és erőforrás-igényesek, továbbá nehezen skálázhatók a többdoménű, többnemzeti műveletek adat- és dokumentumterheléséhez. MI-eszközök alkalmazásával a validáció három szinten gyorsítható: (1) formai és logikai konzisztencia ellenőrzése a tervdokumentumok között; (2) végrehajthatósági és erőforrás-összhang vizsgálata (idő, készletek, mozgások, kapacitások); (3) hatás- és kockázatbecslés szimulációval, illetve prediktív modellekkel. A cél nem a parancsnoki felelősség kiváltása, hanem a döntési ciklus lerövidítése és a hibák korai kiszűrése.

8.1.1 Tervvalidáció: megfelelés, konzisztencia és végrehajthatóság

Tervezési környezetben a validáció gyakran a „józan ész” és a szakértői tapasztalat implicit tesztje. A végrehajtás során azonban a tervet explicit módon kell ellenőrizni: milyen feltételezések sérültek, mely döntési pontok közelednek, és mely erőforrás- vagy időkorlátok válnak kritikus szűk keresztmetszetté. AJP-5 a tervezés során is kiemeli a feltételezések folyamatos felülvizsgálatát; ez a végrehajtási fázisban a monitorozás egyik alapfeladata.[34]

A validáció célszerűen három kategóriába rendezhető:

1. **Megfelelési validáció:** a terv összhangja a politikai-stratégiai iránymutatással, a műveleti mandátummal, a szabályozó dokumentumokkal (pl. ROE), valamint a partner- és befogadó nemzeti korlátozásokkal. Ebben a körben az MI elsősorban dokumentum- és tudásbázis-kereséssel, valamint követelmény-nyomonkövetéssel segíthet.
2. **Konzisztencia-validáció:** az OPLAN és mellékletei (például logisztika, kommunikáció és információs rendszerek, felderítés, célrendszer) közötti logikai ellentmondások felderítése. Itt a természetesnyelv-feldolgozás (NLP) és a tudásgráf-alapú ellenőrzés képes gyorsan kiszűrni az eltérő fogalomhasználatot, a hiányzó hivatkozásokat, vagy a nem kompatibilis ütemezéseket.

3. Végrehajthatósági validáció: annak vizsgálata, hogy az erők, eszközök, idő és készletek ténylegesen rendelkezésre állnak-e a feladatok teljesítéséhez, továbbá a terv időkritikus elemei (mozgások, átcsoportosítás, tűztámogatás, logisztikai utánpótlás) nem ütköznek-e erőforrás- vagy hozzáférési korlátokba. Itt optimalizálási és szimulációs módszerek (pl. diszkrét esemény szimuláció, Monte Carlo, ügynök-alapú modellezés) kombinálhatók MI-algoritmusokkal a gyors „mi lenne ha” elemzésekhez.

Az MI-alapú validáció gyakorlati értéke abban rejlik, hogy a fenti három kategóriában egyaránt képes folyamatos „riasztási” logikát működtetni: nemcsak egyszeri ellenőrzést végez, hanem a beérkező jelentések és szenzoradatok alapján jelzi, ha a terv valamely feltételezése romlik, vagy egy kritikus erőforráspálya (pl. üzemanyag, muníció, egészségügyi kapacitás) a küszöbérték alá esik.

8.1. táblázat – MI által támogatott tervvalidációs ellenőrzési pontok

Validációs fókusz / kérdés	MI-támogatás típusa	Bemeneti adatok / artefaktumok
Feltételezések romlása (assumptions at risk) kimutatható-e?	Anomáliaészlelés, trend- és küszöbfigyelés, bayesi frissítés	Jelentések, ISR-adatok, logisztikai készletszintek, meteorológia
OPLAN mellékletek közötti ellentmondások (idő, tér, feladat) vannak-e?	NLP-alapú követelmény-kivonatolás; tudásgráf; szabályalapú ellenőrzés	OPLAN/CONOPS szövegek, mellékletek, ütemtervek, feladatlisták
Kritikus erőforrások (pl. üzemanyag, EOD, MEDEVAC) kapacitása elegendő-e?	Optimalizálás, szimuláció; erőforrás-útvonal elemzés	Erő- és eszközjegyzék, logisztikai adatok, útvonal/hálózat modell
ROE és politikai korlátozások sérülésének kockázata hol nő?	Szabálybázis + szemantikus keresés; riasztások a feladatokhoz rendelve	ROE, parancsnoki iránymutatás, jogi mellékletek, céljegyzék
Várható ellenséges reakciók és időkritikus döntési pontok előre jelezhetők-e?	Prediktív modellek; játékelméleti/ügynök-alapú szimuláció; COA-értékelés	Ellenségmodell, eseménynaplók, korábbi mintázatok, aktuális helyzetkép
Civil hatások és átmeneti feltételek teljesülése mérhető-e?	NLP a civil jelentésekből; CV GEOINT/IMINT adatokból; indikátor-dashboards	CIMIC adatok, NGO jelentések, OSINT, térképek, képi adatok

A fenti ellenőrzési pontok közös jellemzője, hogy explicit adatmodellezést és metaadatolást igényelnek. A validáció csak akkor automatizálható, ha a terv-elemek (feladatok, erőforrások, időzítés, korlátozások) géppel feldolgozható reprezentációban is elérhetők. Ezért a COPD 5-6 fázisában a törzs digitális „terv-nyomvonalat” (traceability) hoz létre: minden jelentős feladat és döntési pont visszavezethető a parancsnoki szándéokra, a kritikus feltételezésekre és a mérőszámokra (Measures of Performance/Effectiveness), miközben a rendszer naplózza, hogy a terv mikor és miért módosult.

8.1.2 MI-alapú validációs eszköztár: szabályok, szimulációk és „digitális iker” megközelítés

A tervvalidáció MI-eszköztára akkor robusztus, ha több, egymást kiegészítő komponensre épül. A kizárólag neurális modellekre támaszkodó megoldások (pl. nagy nyelvi modellek) gyorsak és rugalmasak, de megbízhatósági és magyarázhatósági kockázatokat hordoznak. Ezzel szemben a szabályalapú rendszerek determinisztikusak, de nehezen tarthatók karban és korlátozottan általánosíthatók. A végrehajtási fázisban ezért jellemzően hibrid megközelítés szükséges, amely a „szabályok + statisztika + szimuláció” hármására támaszkodik.[12], [21]

Szabály- és követelményellenőrzés: A tervdokumentumokból az MI képes követelményeket és korlátozásokat kinyerni (például: milyen erőből, milyen időablakban, milyen területen, milyen együttműködési feltételek mellett kell feladatot végrehajtani). A kinyert elemeket egy tudásgráfba vagy formális ontológiába illesztve automatikus konzisztencia-ellenőrzések futtathatók. Például jelölhető, ha egy alárendelt feladat olyan erőre támaszkodik, amely másutt ugyanabban az időablakban már lekötött, vagy ha a kommunikációs függelék nem tartalmazza a kritikus együttműködési csatornát egy többnemzeti komponens számára.

Szimulációs validáció: A diszkrét esemény szimuláció és az ügynök-alapú modellezés alkalmas arra, hogy a végrehajtási ütemterv „torlódási” pontjait, logisztikai csúcsterheléseit, valamint a várható késéseket láthatóvá tegye. A szimuláció akkor különösen értékes, ha a bemeneti paraméterek bizonytalanok (például útvonalbiztonság, hidak teherbírása, időjárás, ellenséges zavarás). Ilyenkor a Monte Carlo módszerek a kockázati tartományt és a valószínűségi kimeneteket adják meg, amely a parancsnoki döntéshez használható „bizonytalansági keretet” biztosít.

Digitális iker (digital twin) a végrehajtásban: A digitális iker megközelítés a terv és a valóság közötti folyamatos összevetést jelenti. A tervből származó elvárt állapotokat (pl. terület-ellenőrzés, logisztikai készletszintek, beérkezési idők, szolgáltatási szintek) összeveti a valós idejű adatokkal, majd a különbségekből következtet a trendekre és a várható eltérésekre. A NATO digitális transzformációs dokumentumai a data-driven döntéshozatal és a többéves, „ember-folyamat-technológia” alapú modernizáció szükségességét hangsúlyozzák, amely közvetlenül támogatja a digitális iker jellegű megoldások bevezetését műveleti szinten.[58]

„Plan QA” nagy nyelvi modellekkel: A nagy nyelvi modellek (LLM-ek) a tervvalidációban leginkább úgy hasznosíthatók, mint strukturált „kérdézz-felelek” és összefoglaló réteg. Képesek a hosszú mellékletek gyors kivonatolására, az ellentmondások gyanús pontjainak kiemelésére, valamint a döntéshozó számára rövid, kontextusgazdag tájékoztatásra. Ugyanakkor a hallucináció, az adatbiztonság és a forráskövethetőség miatt a LLM-eket csak ellenőrzött környezetben, auditálható adatforrásokra támaszkodva, emberi felülvizsgálattal célszerű alkalmazni. A NIST AI RMF a kockázatok feltérképezését, mérését és kezelését szervezeti szinten keretezi, ami a tervvalidációba beépített GenAI eszközök esetén is releváns.[12]

8.1.3 Adaptív újratervezés: döntési pontok, ágak és folytatások MI-támogatással

Az adaptív újratervezés lényege, hogy a törzs már a tervezési fázisokban azonosítja azokat a döntési pontokat (decision points), amelyeknél a végrehajtás „átfordulhat” másik ágra (branch) vagy folytatásba (sequel). A döntési pontokhoz indikátorokat és küszöbértékeket rendel (például: ellenséges erők átcsoportosítása, kritikus infrastruktúra sérülése, civil támogatottság változása, logisztikai készlet minimum alá csökkenése). A végrehajtás során az információgyűjtés és -feldolgozás feladata az, hogy ezekre az indikátorokra időben „rávilágítson”. [59], [36]

A felderítési és hírszerzési támogatás (intelligence support) doktrinális kereteit a JP 2-01 és a JP 2-0 tárgyalja; ezek rámutatnak, hogy a műveletek teljes életciklusában szükséges a követelmények menedzselése, a gyűjtés-értékelés-szétosztás (intelligence cycle) működtetése, valamint az információk hozzáférhetővé tétele a döntéshozók számára.[59], [38] Az MI itt két módon kapcsolódik: (1) automatizálja a döntési pontokhoz kötött indikátorok folyamatos figyelését; (2) javaslatokat generál a lehetséges korrekciókra, amelyek megfelelnek a terv korlátozásainak.

Az adaptív újratervezés MI-támogatása tipikusan egy „zárt láncú” döntéstámogató rendszeren keresztül valósul meg:

1. Észlelés (detect): esemény- és anomáliaészlelés streaming adatokon (jelentések, ISR, logisztika, kibertér, média).
2. Diagnózis (diagnose): ok-okozati hipotézisek generálása (miért történik az eltérés), és alternatív magyarázatok súlyozása.
3. Opciók (generate options): a rendelkezésre álló erőforrások és korlátozások alapján javasolt módosítások (pl. átcsoportosítás, ütemezés-változtatás, másik hatáslánc), összhangban a parancsnoki szándékkal.
4. Döntés és végrehajtás (decide/act): emberi döntés (human-in-the-loop), majd a döntés visszavezetése a tervadatmodellbe és az assessment keretrendszerbe.

A rendszer megbízhatóságának előfeltétele a kockázatkezelés és a kormányzás. A NIST AI RMF a kockázatok feltérképezését, mérését és kezelését iteratív, életciklus-szemléletű megközelítésben javasolja, amely jól illeszkedik a végrehajtási fázisban jelentkező modelldrift és adatminőségi problémák kezeléséhez.[12]

8.2. táblázat – Adaptív újratervezés tipikus triggerei és MI-támogatás

Trigger / indikátor	Lehetséges műveleti következmény	MI-támogatás	Ajánlott kontroll
Ellenséges erők mintázatváltozása (pl. új tűzeszközök megjelenése)	Védelmi, manőver és erővédelem módosítása	Mintázatfelismerés, prediktív elemzés, szimuláció	Hírszerzési validáció; alternatív hipotézisek
Saját erők veszteségei/üzemképesség romlása	Feladatsorrend és prioritások módosítása	Késleltetés-becslés, optimalizálás erőforrás-korlátokkal	Parancsnoki döntés; ROE/mandátum ellenőrzés
Logisztikai csomópont sérülése vagy útvonalvesztés	Utánpótlási ütem és manőver ütemezés áttervezése	Hálózat- és útvonaloptimalizálás; kockázati szimuláció	Logisztikai törzsvezető jóváhagyása

Kommunikációs zavarás / kibertámadás	C2 késleltetés, COP romlás, alternatív eljárások	Anomáliaészlelés; támadásattribúció-támogatás; reziliencia-javaslat	CIS/CYBER szakértői ellenőrzés; incidenskezelési protokoll
Civil helyzet romlása (tüntetések, humanitárius nyomás)	Legitimitás és átmeneti feltételek veszélye	OSINT/NLP, közösségi média trendek, sentiment	CIMIC és politikai tanácsadó értékelés
Információs térben dezinformációs hullám	Befolyásolás; csapatmorál; partnerségi kockázatok	Deepfake detektálás; hálózatelemzés; narratíva-azonosítás	StratCom/PSYOP S jóváhagyás; forráskritika
Meteorológiai/terepviszonyok váratlan változása	Mozgások és ISR-ciklus átütőmezése	Időjárás- és terepmodellek; késés-forecast	Műveleti vezetés döntése; repülésbiztonság
Politikai korlátozások módosulása (nemzeti caveat)	Feladatok újradelegálása; koalíciós súrlódás	Dokumentum-figyelés; követelmény-nyomonkövetés	Jogi/politikai szintű egyeztetés

Az adaptív újratervezés MI-támogatása a gyakorlatban csak akkor növeli a harcértéket, ha elkerüli az automatizálási torzítást (automation bias) és a „fekete doboz” döntéstámogatást. A DoD felelős MI-implementációs iránymutatása a parancsnoki felelősség és az etikai elvek mentén követeli meg, hogy a döntéstámogató algoritmusok átláthatóak, teszteltek és megfelelően felügyeltek legyenek.[21] E követelmények a NATO műveleti környezetben különösen fontosak a többnemzeti döntési lánc, a nemzeti korlátozások, valamint a jogi megfelelés miatt.

8.2 Valós idejű monitorozás és visszacsatolási hurkok

A COPD 5. fázisának központi terméke a végrehajtás során a megbízható és naprakész helyzetkép (COP) és az operatív értékelés (assessment) fenntartása. A valós idejű monitorozás nem pusztán „adatgyűjtés”, hanem olyan visszacsatolási rendszer, amely a parancsnoki

szándékhoz rendelt mérőszámok alapján jelzi: a művelet a kívánt irányba halad-e, milyen mellékhatások jelennek meg, és hol szükséges beavatkozás vagy tervmódosítás.[34]

A JP 2-0 és a JDP 2-00 egyaránt kiemeli, hogy az intelligencia tevékenység célja a döntéshozók információigényének kielégítése, az információk időben történő feldolgozása és terjesztése, valamint a műveleti tervezés-végrehajtás-értékelés támogatása. A valós idejű monitorozás ezért a hírszerzési lánc és a műveleti vezetés közös, összehangolt feladata.[59], [36]

A modern műveleti környezetben a megfigyelhető jelenségek mennyisége (szenzoradatok, képi hírszerzés, jel- és kibertéri adatok, civil jelentések, média és közösségi platformok) meghaladja a kézi feldolgozás kapacitását. MI nélkül a törzs vagy információs túlterhelésbe kerül, vagy túlságosan lecsökkenti a figyelt indikátorok körét. MI-vel olyan „szűrő és fókuszáló” réteg építhető, amely a döntéshez releváns jelzéseket kiemeli, a tömeges adatokból pedig tömör és visszakövethető tájékoztatót állít elő.

8.2.1 Adatarchitektúra a végrehajtásban: adatfolyamok, metaadatok és döntési ritmus

A valós idejű monitorozás technikai előfeltétele egy olyan adatarchitektúra, amely képes a különböző forrásokból érkező adatokat (szenzor, emberi jelentés, nyílt forrás, műveleti rendszerek logjai) közös metaadatrendszerben kezelni. A NATO digitális transzformációs stratégiai dokumentumai kifejezetten a modern technológiák és az új működési módok kihasználásával kívánják erősíteni a szövetségi döntéshozatalt; ehhez a műveleti szintű adatáramlás egységesítése és az adatok interoperábilis kezelése alapfeltétel.[58]

A praktikus megvalósításban az adatfolyam a következő lépésekre bontható:

1. Ingesztio (ingestion): adatgyűjtés a C2 rendszerekből, ISR-platformokról, logisztikai rendszerekből, valamint nyílt forrásokból. A források minőségének és hozzáférési jogosultságainak rögzítése már itt megtörténik (source tagging).
2. Előfeldolgozás (pre-processing): tisztítás, normalizálás, időbélyeg és georeferencia egységesítése, anonimizálás/PII kezelése, majd a műveleti szintű adatmodellekhez történő illesztés.
3. Elemzés (analytics): streaming analitika (anomáliaészlelés, eseménydetektálás), batch analitika (trendek, predikciók), valamint multimodális feldolgozás (szöveg + kép + jel).

4. Terjesztés (dissemination): a releváns információk eljuttatása a megfelelő döntési szintre (taktikai, műveleti, stratégiai) a „need-to-know” és „need-to-share” elvek kiegyensúlyozásával.[59]

A metaadatok szerepe kiemelt: a forrás megbízhatósága, a begyűjtés módja, a feldolgozási lépések és a bizonytalanság (confidence) jelölése nélkül az MI által előállított riasztások könnyen félrevezetővé válhatnak. A doktrinális elvárások (pl. kontextusban értelmezett információ, alternatív magyarázatok kezelése) csak akkor teljesíthetők, ha az adatokhoz és az MI-kimenetekhez társul a „hogyan jutottunk ide” nyoma (provenance).[60]

8.3. táblázat – Példa műveleti monitorozási területekre, indikátorokra és MI-támogatásra

Terület	Indikátorok (példák)	Adatforrások	MI-technikák	Frissítés
Műveleti előrehaladás (MoP)	feladatok teljesülése, ütemterv-eltérés, késések	C2 eseménynaplók, jelentések	NLP kivonatolás; ütemezés-analitika	perc–óra
Hatás (MoE)	terület-ellenőrzés, ellenség aktivitás, lakossági hozzáállás	ISR, OSINT, CIMIC	predikció; sentiment; CV	óra–nap
Erővédelem	incidensek, fenyegetési szint, konvojkockázat	SIGINT/CYBINT, járőrjelentés	anomália; gráf/hálózat; AML riasztás	perc–óra
Logisztika	készletszintek, fogyási ráta, csúcsterhelések	logisztikai rendszerek, szenzorok	forecast; optimalizálás; szimuláció	óra–nap
Információs környezet	narratívák terjedése, bot-aktivitás, deepfake gyanú	közösségi média, médiafigyelés	NLP topic; hálózatelemzés; deepfake detektálás	perc–óra
Kibertér/CIS	szolgáltatás-kiesés, laterális mozgás jelei, zavarás	SIEM, hálózati logok	anomália; klaszterezés; korreláció	másodperc–perc

Átmeneti feltételek	helyi intézmények működése, szolgáltatások helyreállása	civil jelentések, host nation adatok	NLP; idősor; kockázati index	nap-hét
---------------------	---	--------------------------------------	------------------------------	---------

8.2.2 Visszacsatolási hurok: operatív értékelés és „zárt láncú” döntéstámogatás

A monitorozás akkor értékes, ha a beérkező jelekből következtetés születik, majd a következtetés beépül a döntésekbe. Ez a logika írható le „visszacsatolási hurokként”: (1) mérés - (2) értelmezés - (3) döntés - (4) cselekvés - (5) újramérés. A JP 2-0 szemlélete szerint az intelligencia tevékenység végső célja, hogy a műveletekhez releváns, időben elérhető és megfelelően értékelt információt biztosítson; ebben a visszacsatolás kulcsszerepet játszik.[59]

MI alkalmazásával a visszacsatolási hurok két ponton gyorsítható: egyrészt a „mérés” és „értékelés” közötti átmenetben (automatikus eseménydetektálás, indikátor-számítás, bizonytalanság-kezelés), másrészt az „értékelés” és „döntés” között (alternatívák generálása, kockázati rangsor, következmény-szimuláció). Fontos ugyanakkor, hogy az MI ne homogenizálja a gondolkodást: több, egymással versengő hipotézist kell fenntartani, és a bizonytalanságot explicit módon jelezni.[60]

Az ICD 203 analitikai standardjai - többek között - megkövetelik az objektivitást, a források és módszerek transzparenciáját, az alternatív hipotézisek mérlegelését, valamint a bizonytalanság megfelelő kommunikálását. Ezek a követelmények jól illeszkednek az MI-alapú döntéstámogatás tervezéséhez: a rendszernek nemcsak „eredményt” kell adnia, hanem meg kell mutatnia a bizonyítékok és feltételezések struktúráját is.[60]

A visszacsatolás implementálásában gyakorlati mintát ad a „KPI-fa” vagy „hatáslánc” (effects chain) reprezentáció: a parancsnoki célokat alacsonyabb szintű indikátorokra bontja, majd az indikátorokhoz adatforrásokat és frissítési szabályokat rendel. Az MI ebben a modellben a fa alsó szintjein (adatfeldolgozás) és a középső szinteken (mintázatok, predikciók) a legerősebb, míg a felső szinten (célok és elfogadható kockázat) az emberi döntés domináns.

8.2.3 Dezinformáció, deception és adversarial ML a végrehajtási adatokban

A végrehajtási fázisban a „valós idejű” adat nem egyenlő a „valós” adattal. A műveleti környezet szereplői - beleértve az ellenséget - aktívan törekedhetnek arra, hogy a döntéshozatal inputjait torzítsák (deception), vagy a közvéleményt, a helyi lakosságot, illetve a koalíciós kohéziót befolyásolják (dezinformáció).

A JDP 2-00 a hírszerzés, elhárítás és biztonság (counter-intelligence és security) integrált szerepét hangsúlyozza a joint műveletek támogatásában; ennek része a megbízhatóság és a biztonsági kockázatok kezelése a teljes információs láncban.[36]

A generatív MI (különösen a szintetikus média) csökkentheti a rosszindulatú szereplők belépési küszöbét: könnyebbé válik nagy mennyiségű, célzott üzenet és manipulált vizuális tartalom előállítása és terjesztése. A mesterséges intelligencia által támogatott dezinformációval foglalkozó friss szakirodalom rámutat, hogy a deepfake tartalmak és a bot-ökoszisztémák egyszerre növelhetik a terjesztés sebességét és a hihetőség látszatát, miközben a befogadói reakciókat is erősíthetik.[39]

A műveleti döntéstámogató rendszerek ezért nem tekinthetnek adottnak a bejövő információk integritását. A valós idejű monitorozás MI-alkalmazásainál célszerű bevezetni egy „bizalomréteget” (trust layer), amely minden adatponthoz megbízhatósági és manipulációs kockázati jelölést társít. A trust layer tipikus összetevői:

- Forrásmegbízhatóság és proveniencia: az adat eredetének, gyűjtési módjának és feldolgozási láncának nyomon követése (audit trail).
- Tartalomelemzés: NLP és multimodális modellek a narratívák, anomáliák, koordinált viselkedés és deepfake gyanú jelzésére.
- Kereszt-validáció: több, független adatforrás közötti konzisztencia vizsgálata (például képi és jelhírszerzés, helyszíni jelentések, OSINT).
- Adversarial ML védelem: a modellek ellenálló képességének tesztelése és a támadási felület csökkentése (pl. input-szűrés, robusztus tréning, red teaming).

A trust layer kialakítása a NIST AI RMF életciklus-alapú kockázatkezelési szemléletéhez illeszkedik: a kockázatokat azonosítani (map), mérni (measure) és kezelni (manage) kell, miközben a kormányzás (govern) biztosítja az elszámoltathatóságot.[12]

Fontos korlát, hogy a dezinformáció felismerése sok esetben nem tisztán technikai kérdés: a kontextus, a kulturális kódok és a politikai helyzet döntő tényezők. Ezért az MI által jelzett mintázatokat célszerű összekapcsolni a humán elemzéssel, valamint a kommunikációs és civil-műveleti szakértőkkel (StratCom, CIMIC). A cél a döntéshozatal „szennyeződésének” csökkentése úgy, hogy közben ne blokkoljuk a valós és releváns információ gyors áramlását.

Összességében a valós idejű monitorozás MI-támogatása akkor tekinthető érettnek, ha a helyzetkép és az értékelés nem csupán adatvizualizáció, hanem a döntési ritmushoz illesztett, bizonytalanságot kommunikáló és manipulációval szemben reziliens visszacsatolási rendszer. Ez a képesség közvetlenül befolyásolja a 6. fázis sikerét is, mivel az átmenet feltételei csak megbízható, mérhető és auditálható indikátorok alapján értékelhetők.

A végrehajtás során a műveleti környezetre vonatkozó előfeltevések gyorsan változhatnak; ezért a közös hírszerzési előkészítés (JIPOE) termékeit célszerű „élő dokumentumként” kezelni. A JP 2-01.3 szerint a JIPOE célja, hogy a parancsnokot és a törzset az ellenség és a környezet várható magatartásának megértésével támogassa; a végrehajtási monitorozás MI-eszközei (anomáliaészlelés, predikció, vizuális analitika) így közvetlenül hozzájárulhatnak a JIPOE frissítéséhez és az aktuális helyzetértékeléshez.[10]

8.3 Művelet utáni elemzés MI segítségével

A COPD 6. fázisában (Transition) a műveleti tevékenységek átadása vagy lezárása mellett a tudásmenedzsment és a tanulságok rendszerezése is kiemelt szerepet kap. A művelet utáni elemzés (after-action review, AAR) célja nem a „múlt dokumentálása”, hanem a következő műveletek és képességfejlesztés támogatása: a sikeres minták azonosítása, a hiányosságok okainak feltárása, valamint a doktrinális és szervezeti tanulságok rögzítése.

A művelet utáni elemzés jellegzetesen több, heterogén adathalmazt integrál: eseménynaplókat és parancsnoki jelentéseket, felderítési termékeket, logisztikai és karbantartási adatokat, kommunikációs és kibertéri naplókat, valamint civil és partneri visszajelzéseket. Emberi erőforrással mindez csak korlátozott mélységben dolgozható fel; ezért a MI itt az „adatbányászat” és a „tudáskinyerés” nagy részét automatizálhatja, miközben az értelmezés és a következtetés továbbra is emberi felelősség marad.[61]

A művelet utáni elemzésben a MI három tipikus értékláncban jelenik meg: (1) automatizált adatösszerendezés és dokumentumfeldolgozás; (2) mintázatok és ok-okozati összefüggések feltárása; (3) tanulságok és ajánlások strukturált formában történő rögzítése és visszacsatolása a képzésbe, doktrínába és a jövőbeli tervezésbe.

8.3.1 AAR adatforrások, adatminőség és bizonytalanság

Az AAR-hoz rendelkezésre álló adatok minősége egyenlőtlen: sok forrás strukturálatlan (szöveges jelentések), mások hiányosak (szenzorok kiesése), vagy eltérő osztályozási szinten keletkeznek. A MI rendszerek számára ez egyszerre jelent lehetőséget és kockázatot: a feldolgozás gyorsítható, de a hibás, hiányos vagy torz adatokból téves tanulságok születhetnek (garbage in - garbage out).

Az ICD 203 analitikai standardjai a források és módszerek transzparens kezelését, a bizonytalanság jelzését és a következtetések alátámasztását követelik meg; ezek az elvek az AAR MI-támogatásában is közvetlenül alkalmazhatók.[60] A gyakorlatban ez azt jelenti, hogy a MI-eszköznek minden kivonat, klaszter vagy „tanulság-javaslat” mellé meg kell adnia az evidenciák listáját (mely dokumentumok, események, adatsorok támasztják alá), és a bizalom szintjét.

A földi igazság (ground truth) kialakítása különösen nehéz összhaderőnemi műveletekben. Az események több forrásban, eltérő időbélyeggel és eltérő értelmezési kerettel jelennek meg. Ilyenkor a cél nem egy „tökéletes” igazság, hanem egy verifikálható esemény- és hatásmodell létrehozása: mely eseményeket tekintünk alapvető ténynek (confirmed), melyek hipotézisnek (assessed), és melyek megkérdőjelezettnek (disputed). A kategorizálás a későbbi MI-tréning és -értékelés előfeltétele.

A NIST AI RMF a mérés és menedzsment szakaszban kiemeli a teljesítménymutatók és a kockázati mutatók együttes kezelését. AAR-környezetben ez például azt jelenti, hogy az automatikus szöveg-osztályozó modell pontossága önmagában nem elég; mérni kell a hamis pozitívak (téves tanulságok) és hamis negatívok (elmulasztott tanulságok) műveleti következményeit is.[12]

8.3.2 Automatizált feldolgozás: NLP, multimodális elemzés és hálózati megközelítés

Az AAR automatizálásának „belépési pontja” tipikusan a dokumentum- és eseménybányászat. A jelentésekből (SITREP, INTREP, logisztikai jelentések) NLP-vel kinyerhetők a fő entitások (személyek, alegységek, helyek, objektumok), események (pl. kontakt, támadás, késedelem, technikai hiba) és a hozzájuk tartozó idő- és térparaméterek. Ezzel létrejön egy géppel feldolgozható eseménygráf, amelyre már futtathatók hálózati és időbeli elemzések.

A multimodális feldolgozás kiterjeszti a képet: képi adatokból (GEOINT/IMINT) számítógépes látás (CV) segítségével detektálhatók infrastruktúra-változások, mozgások, sérülések vagy objektumok. A képi és szöveges információ integrálása különösen hasznos a civil hatások és az átmeneti feltételek vizsgálatában (például: helyreállt-e egy kritikus útvonal, működik-e egy szolgáltatás).

A hálózatelemzés (gráf- és közösségetektálás) az AAR-ban két módon értékes: egyrészt feltárja a kulcsszereplőket és kapcsolatokat (például: információ-áramlási csomópontok, koordinációs minták), másrészt alkalmas a rendszerhibák gyökérokainak (root cause) vizsgálatára, ha a hibák láncolata egy gráfban követhető. Ilyen lehet például egy ellátási láncban megjelenő késés, amely több alegységnél okoz dominóhatást.

A nagy nyelvi modellek a művelet utáni elemzésben különösen a „szöveg->struktúra” és a „struktúra->szöveg” átmenetekben hasznosak: képesek nagy mennyiségű dokumentumot összefoglalni, tematikusan csoportosítani, majd a kivonatokat egységes szerkezetű AAR jelentéssé formálni. Ezzel párhuzamosan azonban a forráskritika és a visszakereshetőség követelménye (mely állítás mely dokumentumon alapul) továbbra is kulcstényező, különösen többnemzeti környezetben.[60]

A NATO STO S&T Trends jelentése a 2020-2040-es időszakra a technológiák egyre inkább intelligenssé és összekapcsolttá válását emeli ki, ami az AAR folyamatát is átalakítja: a jövőben várhatóan nagyobb hangsúlyt kap az automatizált adatgyűjtés, a szenzorhálózatok és a digitális rendszerek naplóinak elemzése, valamint a gyors tudásmegosztás.[62]

8.4. táblázat – Művelet utáni elemzési feladatok és MI-eszközök hozzárendelése

AAR feladat	Bemenet	MI-technika	Kimenet
Idővonal rekonstruálása	eseménynaplók, jelentések, szenzoradatok	eseménykivonatolás (NLP), idősor-illesztés	hitelesített timeline, kritikus pontok
Sikerek és hibák klaszterezése	AAR jegyzetek, incidensleírások	topic modeling, klaszterezés	ismétlődő minták, témacsoportok
Gyökérok-elemzés	folyamatlépések, logisztikai és CIS logok	gráf-okozatiság, korrelációs elemzés	okok rangsora, hozzájárulási arányok
Hatásvizsgálat	MoE/MoP indikátorok, civil adatok	predikció, kontrafaktuális elemzés (korlátozottan)	hatáslánc értékelés, kockázatok
Lessons learned kivonatolás	több száz oldal jelentés	LLM alapú összefoglalás + forráshivatkozás	strukturált tanulságlista
Képzési és doktrinális javaslatok	tanulságok és teljesítménymutatók	szabály- és sablonalapú generálás, szakértői review	frissített SOP/doktrína javaslatok

8.3.3 Tudásmenedzsment és átmenet: a tanulságok visszacsatolása a következő ciklusba

A Transition fázis egyik kockázata, hogy a művelet végén a tudás „szétszóródik”: a kulcsemberek rotálnak, a rendszerlogok archiválódnak, a partneri szervezetek eltérő formátumban adják át a tapasztalatokat. MI-vel kialakítható egy olyan szervezeti tudástár, amelyben a tanulságok nem csupán narratív leírások, hanem kereshető, címkézett és összekapcsolt elemek (például: „logisztika - készletgazdálkodás - X típusú hiány”; „CIS - zavarás - redundancia”).

Ebben a modellben a tudásgráf központi szerepet kap: összekapcsolja az eseményeket, döntéseket, erőforrásokat és hatásokat. Így a következő tervezési ciklusban a törzs nem „nulláról” indul, hanem azonnal lekérdezheti: adott környezetben (PMESII) milyen kockázatok és bevált megoldások jelentek meg korábban. A JP 2-0 a „intelligence lessons learned” elvét is hangsúlyozza; a MI ezt a tanulási folyamatot skálázhatóbbá teheti.[59]

A tudásmenedzsment ugyanakkor nem csak technológia. Jensen hangsúlyozza, hogy a jövő intelligenciaelemzése a technológia, az emberek és a folyamatok együttes fejlesztését igényli; a tanulságok akkor hasznosulnak, ha intézményesült folyamat (standardizált AAR, felelősök, visszamérés) biztosítja a beépülést.[61]

Clapper memoárja (a hírszerzési ökoszisztéma belső logikáját bemutatva) rámutat arra is, hogy a modern információs környezetben a transzparencia, a hitelesség és a bizalom kérdései felértékelődnek. A tanulságok kommunikálásánál ezért fontos a világos állítás-bizonyíték logika és a következtetések korlátainak jelzése, különösen többnemzeti döntéshozatal esetén.[63]

8.4 Megvalósítási kihívások (pl. adatbiztonság, interoperabilitás)

A COPD 5-6 fázisában alkalmazott MI-rendszerek bevezetése nem pusztán technikai fejlesztés: egyszerre érint adatkezelési, biztonsági, interoperabilitási, jogi-etikai, szervezeti és képzési kérdéseket. A végrehajtási fázis időnyomása miatt ezek a kihívások különösen élesen jelentkeznek: ami a tervezés során „kényelmi funkció”, az a végrehajtásban döntéskritikus rendszerkomponenssé válik. Az alábbiakban a legfontosabb kockázatokat és kezelési mintákat foglaljuk össze a nemzetközi keretek (NIST AI RMF, DoD Responsible AI, EU AI Act) és a szövetségi digitális törekvések alapján.[12], [21], [64], [58]

8.4.1 Adatbiztonság, kiberreziliencia és ellátási lánc

A végrehajtási fázisban működő MI-rendszer egyszerre adatfogyasztó és -előállító: a bemenetek (szenzorok, jelentések) gyakran minősített környezetből származnak, a kimenetek pedig döntéskritikus információk. Ezért az adatbiztonság két dimenzióban értelmezendő: (1) a klasszikus információbiztonság (bizalmasság, sértetlenség, rendelkezésre állás - CIA triád), valamint (2) a modellbiztonság (adversarial ML, adatfertőzés, modell-exfiltráció).

A NIST AI RMF külön kiemeli, hogy az AI-kockázatok részben eltérnek a hagyományos szoftverkockázatoktól, mert a rendszer teljesítménye az adateloszlás változásától (drift) és az ellenség aktív beavatkozásától is függ.[12] Műveleti környezetben tipikus fenyegetés a „data poisoning”: az ellenség úgy torzítja a beérkező adatokat, hogy a modell téves riasztásokat

adjon, vagy ne jelezzen valódi veszélyt. Ezzel rokon kockázat az input-manipuláció (adversarial examples) képi és jeladatok esetén.

A DoD felelős MI-implementációs dokumentuma a megfelelő tesztelést, ellenőrzést és folyamatos felügyeletet (monitoring) hangsúlyozza. Ennek részeként a műveleti MI-képességekhez célszerű „biztonsági minimál” követelményeket rögzíteni: modellverziók nyomonkövetése, aláírt modellesomagok, naplózás, jogosultságkezelés, valamint rendszeres red teaming és sebezhetőségi vizsgálat.[21]

A többnemzeti műveletekben gyakori a többosztályú (multi-classification) adatkezelés és a cross-domain megoldások igénye. Az MI-alkalmazásoknál ez azt jelenti, hogy a tanítóadatok, a futtatási környezet és a kimenetek terjesztése különböző biztonsági tartományokhoz kötődhet. A cél ezért a „biztonságos adatszétválasztás” és a minimális adat-megosztás elve: csak a szükséges jellemzők (features) kerüljenek át, megfelelő anonimizálással és auditálható szabályokkal.

Az ellátási lánc biztonsága a modellek és könyvtárak eredete miatt szintén kritikus. A műveleti MI rendszerek gyakran nyílt forrású komponensekre épülnek; ezek frissítése, licencelése és sérülékenysége műveleti kockázat. Ezért a bevezetés része kell legyen egy „ModelOps/MLOps” fegyelem: verziókezelés, reprodukálható tréning, konfigurációmenedzsment, és a terepi környezetben (edge) történő biztonságos frissítés.

8.4.2 Interoperabilitás és adatstandardok a szövetségi végrehajtásban

A COPD 5-6 fázisában a döntéstámogatás minősége nagyban függ az interoperabilitástól: a különböző nemzeti rendszerekből és szenzorokból származó adatok csak akkor fuzionálhatók, ha közös vagy átjárható adatmodellek és metaadat-sémák állnak rendelkezésre. A NATO digitális transzformációs stratégiája kifejezetten a működés átalakítását célozza modern technológiák kihasználásával; ennek gyakorlati leképeződése a műveleti szinten az adatintegráció és a közös adatszótár megteremtése.[58]

Az interoperabilitás MI-környezetben legalább négy rétegben értelmezhető:

4. Technikai interoperabilitás: hálózati kapcsolódás, protokollok, adatátvitel és hozzáférés, megbízható időszinkron.
5. Szintaktikai interoperabilitás: az adatformátumok és struktúrák egységesítése (mezők, egységek, időbélyeg, geokód).

6. Szemantikai interoperabilitás: az adatelemek jelentésének egyeztetése (mit jelent egy „incidens”, milyen kategóriák vannak, hogyan értelmezzük a „fenyegetési szintet”).

7. Szervezeti interoperabilitás: a munkafolyamatok, felelősségek és adatmegosztási szabályok összehangolása (need-to-share vs. need-to-know).

AJP-5 és a kapcsolódó tervezési doktrína a többnemzeti környezetben külön hangsúlyt fektet a közös terminológia és a koordináció jelentőségére. MI-eszközök bevezetésekor ez a követelmény megsokszorozódik, mivel a modellek csak akkor adnak konzisztens eredményt, ha a bemenetek összehasonlíthatók.[34]

A szemantikai interoperabilitás megoldására gyakorlati módszer a közös ontológia vagy tudásgráf, amely a kulcsfogalmakat (erők, képességek, helyek, események, hatások) egységesen definiálja, és képes kezelni a nemzeti eltéréseket (mapping). Ez a megközelítés a végrehajtás során a gyors integrációt és a visszakereshetőséget támogatja: ha egy szövetséges partner eltérő kategóriákat használ, a rendszer mégis képes közös dashboardban megjeleníteni az indikátorokat.

A 6. fázis átmeneti feladatai (átadás a helyi intézményeknek, átmenet másik műveleti szakaszba, vagy a művelet lezárása) különösen érzékenyek az interoperabilitásra: az átadási csomagban (handover package) olyan adatoknak és tanulságoknak kell szerepelnie, amelyeket a következő szervezet vagy civil partner is értelmezni tud. Itt az MI segíthet automatizált adat- és dokumentum-konverzióval, illetve a többnyelvű, egységes szerkezetű összefoglalók előállításával.

8.5. táblázat – Interoperabilitási rétegek és tipikus megoldási eszközök MI-környezetben

Réteg	Tipikus megoldások	Megjegyzés a COPD 5-6 fázisban
Technikai	biztonságos hálózati összeköttetés, idősinkron, edge gateway	a valós idejű riasztások késleltetése műveleti kockázat
Szintaktikai	közös adatcsere-formátumok, egységes egységek, ID-k, geokód	a dashboardok és a szimulációk csak tiszta adaton működnek
Szemantikai	ontológia/tudásgráf, fogalomtár, mapping szabályok	többnemzeti eltérések kezelése; visszakereshetőség

Szervezeti	adatmegosztási felelősségi mátrix, SOP, battle rhythm	átmenetnél (handover) a folyamatok összehangolása döntő
------------	---	---

8.4.3 Kormányzás, felelős MI és jogi megfelelés

A kormányzás (governance) a műveleti MI egyik legnagyobb kihívása, mert a végrehajtásban az algoritmusok gyorsan „operatív igazsággá” válhatnak: a dashboardon megjelenő riasztás vagy előrejelzés befolyásolja a parancsnoki döntést, még akkor is, ha a modell bizonytalan. A NIST AI RMF ezért a kockázatkezelés fölé helyezi a kormányzási funkciót (govern): egyértelmű felelőségek, felügyeleti mechanizmusok, dokumentálás és folyamatos értékelés nélkül a technikai teljesítmény nem garantálja a megbízható működést.[12]

A DoD „Implementing Responsible Artificial Intelligence” iránymutatása a felelős alkalmazás elveit (pl. governance, tesztelés, monitoring és emberi felügyelet) a védelmi szervezetek sajátosságaira alkalmazza. A dokumentum lényegi üzenete, hogy az MI rendszerek bevezetésekor már a fejlesztés korai szakaszában rögzíteni kell a célokat, a használati korlátokat, a mérőszámokat, valamint a felelősségi és jóváhagyási láncot.[21] A COPD 5-6 fázisában ez különösen fontos, mert a modellek által generált javaslatok a végrehajtásban azonnali erőalkalmazási vagy erőforrás-allokációs következményekkel járhatnak.

Az Európai Unió AI Act (Regulation (EU) 2024/1689) harmonizált szabályrendszert hoz létre a kockázat-alapú megközelítés mentén. Ugyanakkor a rendelet kifejezetten rögzíti, hogy a kizárólag katonai, védelmi vagy nemzetbiztonsági célokra piacra helyezett, üzembe helyezett vagy használt MI rendszerek kívül esnek a hatályán; amennyiben azonban egy ilyen rendszer más célra (például civil vagy humanitárius) kerül alkalmazásra, a rendelet követelményei relevánssá válhatnak.[64] Ez a kettősség a NATO műveleti környezetben gyakori: ugyanaz a technológia (pl. képfeldolgozó modell) katonai és civil célú helyreállítási feladatot is támogathat.

A jogi megfelelés mellett a „felelős MI” kérdése szövetségi legitimációs tényező: a NATO Strategic Concept a demokratikus értékek és a társadalmi ellenálló képesség fontosságát hangsúlyozza; ezzel összhangban a műveleti MI-alkalmazásoknak támogatniuk kell a transzparenciát, a számonkérhetőséget és az arányosság elvét.[57]

A gyakorlatban a kormányzás kézzelfogható artefaktumokban jelenik meg: kockázati regiszter (AI risk register), modellkártyák (model cards), adatlapok (datasheets), tesztjegyzőkönyvek, és

- különösen generatív MI esetén - használati szabályok (prompt policy), naplózás és audit. Ezek az elemek nem „adminisztratív terhek”, hanem a parancsnoki döntés védőkorlátai: lehetővé teszik, hogy a rendszer használata utólag visszakövethető és értékelhető legyen.

8.4.4 Szervezeti, képzési és bevezetési kihívások: képességfejlesztés a Tranzíciós fázis tükrében

A COPD 5-6 fázisában alkalmazott MI rendszerek sikeres bevezetése a szervezet „tanulási képességén” múlik. A technológia önmagában nem oldja meg a döntéshozatali problémákat; a folyamatokat és a szerepköröket hozzá kell igazítani az adatalapú működéshez. Jensen megfogalmazása szerint a jövő intelligenciaelemzése nemcsak technológiai, hanem emberi és folyamatok kérdése is.[61]

A NATO Digital Transformation Implementation Strategy a transzformációt nem csupán informatikai modernizációként, hanem a működési módok átalakításaként értelmezi. Műveleti szinten ez például azt jelenti, hogy a battle rhythm-be beépülnek az MI-alapú dashboardok, a modell-riasztások validációs lépései, és az assessment fórumok (pl. Joint Effects Board) döntései géppel követhetően kerülnek vissza a tervadatmodellbe.[58]

A NSCAI Final Report a nemzetbiztonsági alkalmazások esetén hangsúlyozza a tehetség, a szervezeti átalakítás és a partnerségek szerepét. Ezek a javaslatok a NATO és nemzeti keretekben is relevánsak: a MI-képességekhez olyan szakemberek (data engineer, data scientist, MLOps, cyber) kellenek, akik képesek a műveleti törzsszel együtt dolgozni és a megoldásokat gyors iterációban fejleszteni.[65]

A magyar nemzeti környezetben a Nemzeti Katonai Stratégia 2030-ig iránymutatást ad a katonai dimenzióban értelmezett kihívások kezelésére. A többdoménű fenyegetések és a modernizációs célok összhangba hozhatók a felelős MI bevezetésével: a cél az, hogy a Magyar Honvédség a szövetségi műveletekben interoperábilis, biztonságos és mérhető MI-támogatást tudjon alkalmazni, különösen a végrehajtási és átmeneti időszakban, amikor a döntési ciklus a leggyorsabb.[11]

A bevezetési kihívások között gyakran alulértékelt a változásmenedzsment: a törzsek részéről természetes ellenállás jelenik meg, ha az MI „fekete doboz” módon javasol intézkedéseket. A sikeres bevezetés ezért lépcsőzetes (spirális) módszertant igényel: először alacsony kockázatú területeken (pl. dokumentum-összefoglalás, adattisztítás, dashboard automatizálás), majd

fokozatosan magasabb kritikalitású döntéstámogatási funkciókban (predikció, optimalizálás), mindig emberi felügyelet mellett.

Az átmenet (Transition) szempontjából különösen fontos a fenntarthatóság: az MI-rendszer ne csak a művelet idején működjön, hanem az átadáskor (handover) és a lezáráskor is képes legyen adatot szolgáltatni és tanulságokat átadni. Ez a követelmény a fejlesztésben azt jelenti, hogy már a tervezéskor definiálni kell a kimenetek „átadhatóságát” (export, interoperábilis formátumok, többnyelvű kivonatok), valamint a hosszú távú archiválás és hozzáférés szabályait.

8.6. táblázat – MI bevezetési kihívások a COPD 5-6 fázisában: kockázatok és mitigációk

Kihívás	Tipikus kockázat	Mitigáció (minták)	Felelős szereplők
Adatbiztonság / osztályozás	adatkiszivárgás; jogosulatlan hozzáférés; cross-domain hibák	zero trust; hozzáférés-kontroll; audit; minimális adat-megosztás	CIS, Security, Data owner
Adatminőség / drift	téves riasztások; romló pontosság; döntéstámogatás félrevezetése	adatminőség-mérők; folyamatos monitor; újratréning; fallback eljárások	MLOps, elemzők, műveleti törzs
Adversarial manipuláció	data poisoning; deepfake; szándékos torzítás	red teaming; input-szűrés; forrás-keresztvalidáció; trust layer	CI, ISR, Cyber, StratCom
Interoperabilitás	adatok nem illeszthetők; közös kép hiánya; koalíciós súrlódás	közös ontológia; mapping; adatcsere-szabványok; közös SOP	J6/J2, partnerek, adatgazdák
Kormányzás / felelősség	automation bias; audit hiánya; felelősség elmosódása	AI risk register; model cards; emberi felügyelet; döntési napló	Parancsnok, jogi tanácsadó, AI governance
Erőforrás és kompetencia	szakemberhiány; túlterhelt törzs; fenntarthatóság hiánya	képzés; vegyes csapatok; iteratív bevezetés; tudásmenedzsment	HR, képzési szervek, parancsnokság

8.5 Összegzés

A COPD 5-6 fázisaiban a siker kulcsa a gyors, mégis megalapozott döntéshozatal: a végrehajtás során a tervet folyamatosan validálni kell, a helyzetképet és a hatásokat mérni kell, majd szükség esetén adaptív újratervezéssel biztosítani a parancsnoki szándék érvényesülését.[56], [34]

A fejezet bemutatta, hogy az MI a tervvalidációban (szabály- és konzisztencia-ellenőrzés, szimuláció, digitális iker), a valós idejű monitorozásban (streaming analitika, anomáliaészlelés, multimodális feldolgozás), valamint a művelet utáni elemzésben (NLP, hálózatelemzés, tudásmenedzsment) képes a döntési ciklust lerövidíteni és a törzsek munkaterhét csökkenteni. Ugyanakkor a dezinformáció, a deception és az adversarial manipuláció miatt a rendszereknek beépített „bizalomrétegre” és forráskritikára van szükségük.[36], [12], [39]

A bevezetés feltétele a felelős MI-kormányzás, az adatbiztonság és az interoperabilitás: a NIST AI RMF és a DoD felelős MI iránymutatása gyakorlati keretet ad a kockázatok életciklus-szintű kezelésére, míg az EU AI Act jogi környezetben releváns megfelelőségi szempontokat rögzít - különösen akkor, ha a katonai rendszerek civil vagy humanitárius felhasználási módokba átcsúsznak.[12], [21], [64]

A NATO stratégiai és digitális transzformációs irányai alapján a jövő műveleti terében a MI alkalmazása nem opcionális, hanem versenyképességi tényező. A COPD 5-6 fázisaiban a legnagyobb hozzáadott érték a mérhető, auditálható és emberi felügyelet alatt álló döntéstámogató megoldásoktól várható, amelyek képesek a többforrású adatokat a parancsnoki döntésekhez közvetlenül hasznosítható tudássá alakítani.[57], [58], [62]

9. FEJEZET - KOCKÁZATOK, ETIKAI ÉS JOGI MEGFONTOLÁSOK, KIBERBIZTONSÁG

Az MI-alapú döntéstámogatás és automatizáció a COPD teljes életciklusán – a helyzetértékeléstől és fenyegetésértékeléstől a katonai válaszopciók kidolgozásán, a végrehajtáson és az átmeneten át – olyan képesség, amely egyszerre növeli a döntések sebességét és a helyzetkép részletességét, ugyanakkor új kockázati felületeket is nyit. A műveleti környezetben alkalmazott MI nem csupán technológiai komponens: szervezeti folyamatokba ágyazott döntési eszköz, amely hat a felelősségi viszonyokra, az elszámoltathatóságra, a bizalomra és a legitimációra is. A kockázatok egy része klasszikus (adatbiztonság, rendszermegbízhatóság), más része kifejezetten MI-specifikus (torzítások, magyarázhatóság, adversarial támadások, generatív „hallucinációk”), és mindezek egy ellentevékenységgel terhelt, gyorsan változó környezetben jelennek meg.[12], [66]

Az utóbbi években több, egymást részben átfedő policy- és standard-keret alakult ki a felelős és biztonságos MI-használat támogatására. Az Egyesült Államok Védelmi Minisztériuma (DoD) a Responsible AI (RAI) bevezetését irányítási (governance) és életciklus-szintű feladatként írja le, hangsúlyozva az etikai elvek operationalizálását, a tesztelés-értékelés (T&E) fontosságát, valamint a privacy és civil liberties védelmét.[21] A NATO AI Strategy a szövetségi interoperabilitás és a közös bizalom szempontjából rögzíti a felelős használat elveit, kiemelve, hogy a képességeket a nemzetközi joggal és a szövetségi értékekkel összhangban kell fejleszteni és alkalmazni.[67]

A NIST AI Risk Management Framework (AI RMF 1.0) egy általános, szervezeti és technikai gyakorlatokat összekapcsoló, életciklus-szemléletű kockázatkezelési modellt ad, amely a kormányzás (Govern), a kontextus-feltérképezés (Map), a mérés (Measure) és a menedzsment (Manage) funkciók köré rendezi az MI-kockázatok kezelését.[12] A generatív MI sajátosságaira a NIST AI 600-1 (Generative AI Profile) külön kockázati leírásokat és kontroll-javaslatokat fogalmaz meg, például a téves állítások (confabulation), a nem kívánt tartalom, a prompt-injekció, a modell- és adat-kiszivárgás, illetve a tartalom-eredet (provenance) kezelésére.[66]

Az európai jogi környezetben az EU AI Act kockázatalapú szabályozási logikát vezet be a tilalmazott és magas kockázatú MI-rendszerek kategóriáival, valamint hangsúlyos

követelményekkel a dokumentáció, naplózás, emberi felügyelet, pontosság/robosztusság és kiberbiztonság terén.[64] Bár a katonai és nemzetbiztonsági alkalmazások kezelése a tagállami hatáskörök miatt sajátos, a rendelet követelménystruktúrája – különösen a bizonyítható megfelelés (auditability) – olyan „best practice” hivatkozási pontot ad, amely a védelmi beszállítói láncban és a dual-use technológiákban közvetetten is megjelenik.[64]

E fejezet célja, hogy a kockázatokat a COPD-hez illesztve, rendszerszemléletben tárgyalja: a felelős MI és emberi kontroll elveit (9.1), a jogi/policy környezetet (9.2), az adat-, torzítás-, magyarázhatósági és információs manipulációs kockázatokat (9.3), az adversarial ML és modellmérgezés elleni védekezést (9.4), végül a minőségbiztosítás és T&E, megfelelés és audit gyakorlati megközelítését (9.5).

9.1. Felelős MI (Responsible AI) és emberi kontroll elvei

A felelős MI katonai kontextusban nem „díszítő etika”, hanem olyan irányítási és mérnöki gyakorlatok rendszere, amely igazolható módon biztosítja: (1) a jogszerűséget és a normákkal való összhangot; (2) a küldetéshez mért, tudatos kockázatvállalást; (3) a hibák és nem várt következmények kezelhetőségét; valamint (4) a döntések visszavezethetőségét és ellenőrizhetőségét. A DoD RAI dokumentum hangsúlyozza, hogy a felelős viselkedésnek a teljes életciklusban – tervezés, fejlesztés, tesztelés, beszerzés, telepítés, üzemeltetés és kivonás – meg kell jelennie, és a szervezetnek képesnek kell lennie bizonyítani, hogy a megfelelő kontrollok működnek.[21]

A NATO AI Strategy a felelős használat elveit a szövetségi interoperabilitás és a bizalom fenntartása érdekében rögzíti. A stratégia alapján a felelős használat nem csupán technikai követelmény, hanem a szövetségi értékek, normák és a nemzetközi jog gyakorlati leképezése az MI-képességek életciklusában.[67] A katonai műveletekben ez különösen ott válik kézzelfoghatóvá, ahol a rendszer kimenete közvetlenül befolyásolhatja a célkijelölést, a tüztámogatás időzítését, a mozgások ütemezését vagy a civil kockázat minimalizálását. E döntési helyzetekben a hibák társadalmi és jogi költsége lényegesen magasabb, mint számos polgári felhasználásban.

A felelős MI egyik központi eleme az emberi kontroll (human oversight/human control). A gyakorlatban ez nem egyszerűen a „human-in-the-loop” jelenlétét jelenti, hanem azt, hogy a döntési felelősség, a jogosultságok és a döntési pontok explicit módon vannak meghatározva.

Célszerű elkülöníteni a felügyeleti módokat: (a) ember a döntési láncban – explicit jóváhagyás szükséges; (b) ember a döntési lánc felett – felülbírási/jóváhagyási lehetőség; (c) ember a rendszer tervezésében – korlátok, tiltások, küszöbök és vészleállítási lehetőségek; (d) ember a „küldetés-kormányzásban” – a rendszer szerepének, céljának, feladatkörének és elfogadható kockázatszintjének meghatározása.[12], [21], [67]

A felelős emberi kontroll nem pusztán szervezeti és jogi kérdés, hanem kognitív és ergonómiai is. Ha a rendszer olyan komplexitású kimenetet ad, amelyet a kezelő nem tud megérteni vagy érdemben mérlegelni (például túl rövid döntési idő, túl nagy adatmennyiség, magyarázat hiánya), akkor a „jóváhagyás” könnyen formalitássá válhat. A NIST AI RMF ezért a kormányzás és mérés funkcióiban olyan bizonyítékokat és metrikákat javasol, amelyek igazolják: a szervezet képes a kockázatokat a valós működésben is felismerni, mérni és kezelni.[12]

A generatív MI megjelenésével a felügyelet új dimenziót kap: a modell nemcsak „prediktál”, hanem szöveget, képet, érvelést, sőt tervvariánsokat generál. Ez egyszerre erősítheti a törzs produktivitását (gyorsabb összegzés, alternatívák), és növeli a félrevezetés kockázatát (plauzibilis, de téves állítások; „hallucinációk”). A NIST AI 600-1 profil a generatív rendszerekre jellemző kockázatokat – például a confabulation/hallucination jelenséget, a prompt-injekciót, a nem kívánt tartalmat és a tartalom-eredet hiányát – külön kontrollcsoportokkal tárgyalja.[66]

Az etikai megfontolások gyakorlati kezelésekor hasznos különválasztani az (a) elv-alapú megközelítést (pl. jogszerűség, emberi méltóság, elszámoltathatóság) és (b) következmény-alapú megközelítést (pl. civil kár minimalizálása, küldetés sikeressége). A modern robotetikai szakirodalom rámutat, hogy az autonóm rendszerek nem válnak „erkölcsi ágensekké” a klasszikus értelemben: a morális felelősség továbbra is emberi és intézményi szinten értelmezendő. Ugyanakkor az automatizáció fokozódása „felelősségi hézagot” teremthet, ha a szerepkörök, döntési pontok és felülvizsgálati mechanizmusok nincsenek világosan kijelölve. A katonai MI-rendszereknél ezért az etikai elemzésnek kézzelfogható kontrollokhoz és döntési szabályokhoz kell kapcsolódnia (felülbírálnak, dokumentált kockázatvállalás, utólagos vizsgálhatóság), különben az etikai elvek nem válnak a műveleti gyakorlat részévé.[68]

9.1.1. RAI elvek operationalizálása a katonai MI-életciklusban

A DoD RAI keretben megjelenő etikai elvek (felelősség, méltányosság, nyomonkövethetőség, megbízhatóság, kormányozhatóság) akkor válnak a gyakorlatban ellenőrizhetővé, ha minden elvhez hozzárendelünk (1) mérhető követelményeket, (2) szervezeti és technikai kontrollokat, valamint (3) bizonyítékokat. A katonai fejlesztésekben ez tipikusan azt jelenti, hogy a képesség-követelményeket (CONOPS/OPLAN mellékletek, SOP-k) összekötjük az MI-fejlesztési artefaktumokkal (adatleltár, model card, tesztsomag, naplózás, jogosultsági rend), és a parancsnoki döntési pontokhoz explicit felügyeleti eljárásokat rendelünk.[21], [12]

A NIST AI RMF „Govern–Map–Measure–Manage” logikája jól használható híd a stratégiai elvek és a mérnöki gyakorlat között. A Govern funkció alatt rögzíthető a felelősségi struktúra, a kockázati tolerancia és a döntési jogkörök; a Map funkcióban feltérképezhetők a felhasználási esetek, a misuse-case-ek, a stakeholderek és a környezeti tényezők; a Measure funkcióban kialakíthatók a teljesítmény-, robosztusság- és torzítás-metrikák; míg a Manage alatt működtethető a folyamatos monitorozás, incidenskezelés és frissítési/rollback mechanizmus.[12] A katonai környezetben a keret külön értéke, hogy a kockázatok a küldetés-kontextushoz köthetők: ugyanaz a modell más kockázati kategóriába eshet adminisztratív és célkijelölési felhasználás esetén.

9.1. táblázat – Felelős MI elvek és bizonyítható kontrollok (katonai döntéstámogató MI)

Elv (RAI/Responsible Use)	Mit jelent a gyakorlatban?	Példa kontrollokra	Bizonyító artefaktumok (audit evidence)
Felelősség és elszámoltathatóság	A döntési felelősség nem delegálható a modellnek; kijelölt felelősök és döntési jogosultságok.	RACI-mátrix; jóváhagyási workflow; döntési küszöbök és „stop” pontok.	Döntési naplók; jogosultság- nyilvántartás; incidensjelentések.
Nyomonkövethetőség és magyarázhatóság	A kimenet inputjai, korlátai és bizonytalansága visszakereshető; a felhasználó érti a kockázatot.	Logging; adat- provenance; magyarázó modulok; uncertainty és confidence kijelzés.	Model card; adatleltár; logok; magyarázati riportok.
Megbízhatóság és roboosztusság	Elvárt teljesítmény valós körülmények között, drift és	Stresszteszt; adversarial teszt; redundáns	T&E jegyzőkönyv; roboosztussági metrikák; drift- riasztások.

	ellentevékenységgel mellett.	érzékelés; fallback eljárások.	
Méltányosság és torzításkezelés	Torzítások feltárása és csökkentése; nem kívánt diszkriminatív hatások kezelése.	Bias audit; reprezentatív adatkészlet; fairness küszöbök; emberi felülvizsgálat.	Bias riport; adatlappal bővített dataset; döntési indoklások.
Kormányozhatóság (governability)	A rendszer korlátozható, leállítható; üzemmódja szabályozható; „biztonsági korlátok” léteznek.	Kill switch; policy engine; szerepkör- alapú hozzáférés; üzemmódváltás.	Konfigurációs jegyzék; hozzáférési log; vészleállítási SOP.

A kockázatalapú arányosítás gyakorlati következménye, hogy a magasabb kockázatú döntésekhez (pl. erőalkalmazás, célprioritás, civil kockázat) szigorúbb kontroll- és bizonyítékcsoport tartozik, míg alacsonyabb kockázatú támogató feladatoknál (pl. dokumentum-összegzés, fordítás, adminisztratív ellenőrzés) a kontrollok arányosíthatók. Ez az elv összhangban áll az AI RMF rugalmas, use-case agnosztikus megközelítésével, és az EU AI Act kockázatalapú logikájával is.[12], [64]

9.1.2. Emberi kontroll a döntési ciklusban: automatizációs torzítás és kalibrált bizalom

A katonai döntéshozatalban az egyik leggyakoribb emberi tényező kockázat az automatizációs torzítás: a kezelők hajlamosak a rendszer által adott javaslatot implicit „hitelesebbnek” tekinteni, különösen időnyomás, információátterhelés és többcsatornás adathalmazok mellett. A felelős MI ezért nem csak a modell teljesítményéről szól, hanem arról is, hogyan prezentáljuk a bizonytalanságot, milyen alternatívákat kínálunk, és hogyan kényszerítjük ki a kritikus ellenőrzési pontokat. A DoD és a NATO dokumentumai a felügyeletet úgy értelmezik, mint a jogszerű és felelős alkalmazás garanciáját, amelynek működnie kell a valós műveleti tempóban is.[21], [67]

A kalibrált bizalom azt jelenti, hogy a felhasználó „pont annyira” bíz a rendszerben, amennyire az bizonyítékok alapján indokolt. Ennek eszközei lehetnek: a kimenetekhez rendelt megbízhatósági szint; a hibamódok ismertetése; a felhasználási korlátok explicit megjelenítése; valamint a „piros zászlók” (red flags) jelzése, amikor a modell a tanítási tartományán kívül

kerül. Az AI RMF Measure–Manage funkciói az ilyen jelzéseket és a monitorozást a kockázatkezelés szerves részének tekintik.[12]

A generatív rendszerekben a túlzott bizalom különösen veszélyes, mert a nyelvi folyékonyság könnyen kompetencia-illúziót kelthet. A NIST AI 600-1 ezért javasolja többek között a forrásmegjelölés és a hivatkozási lánc kényszerítését, a tényállítások gépi és emberi ellenőrzését kritikus feladatoknál, valamint a nem kívánt tartalmak és titokvédelmi kockázatok szűrését.[66]

Az emberi kontroll gyakorlati megvalósítása a törzsmunkában igényli a szerepkörök és döntési pontok formalizálását. Például: a rendszer javasolhat COA-variánsokat, de a kiválasztás és a kockázatvállalás az illetékes parancsnoki szinten marad; vagy a rendszer előkészíthet cél-listákat, de a céljogosultsági és arányossági ellenőrzés többlépcsős, emberi validációhoz kötött. A felelős MI keretek lényege éppen az, hogy ezeket a korlátokat és felelősségi szabályokat a rendszer tervezésébe és üzemeltetésébe is beépítsék.[21], [12]

9.2. Nemzetközi humanitárius jog, NATO/EU/USA policy-k és standardok

A katonai MI-alkalmazások jogi és normatív kerete többretegű: a nemzetközi humanitárius jog (IHL) alapelvei, a szövetségi doktrinák és stratégiák (NATO), a nemzeti policy-k (pl. USA DoD), valamint a regionális jogi környezet (EU) együttesen formálják, hogy egy adott MI-képesség hogyan fejleszthető, telepíthető és használható. A gyakorlati kihívás nem az elvek „ismerete”, hanem azok operationalizálása olyan környezetben, ahol a döntések gyakran valós időben születnek, és az ellenség aktívan törekszik a hibák kiváltására vagy a döntéshozó félrevezetésére.[12], [21], [64], [67]

Az IHL alapelvei (különösen a megkülönböztetés, arányosság és elővigyázatosság) akkor kerülnek fókuszba, amikor a rendszer a „kritikus funkciók” közelébe kerül – azaz a célok kiválasztásának és a támadás végrehajtásának folyamatába közvetlenül beavatkozik. Az ICRC több dokumentuma is arra figyelmeztet, hogy az autonómia növekedése a kritikus funkciókban a human control elvesztésének kockázatát hozza magával, ami jogi és etikai aggályokat vet fel a fegyverhasználat korlátozhatósága és előreláthatósága kapcsán.[69]

A CCW (Convention on Certain Conventional Weapons) keretében működő GGE LAWS folyamat jelzi, hogy az autonóm fegyverrendszerek kérdése továbbra is kiemelt nemzetközi

vita tárgya. A tárgyalások középpontjában a megfelelő korlátozások és a lehetséges szabályozási instrumentumok állnak, valamint az, hogy a meglévő IHL szabályok elegendőek-e, vagy további, specifikus normák szükségesek.[70]

A NATO AI Strategy a felelős használat elveit rögzíti, és az elvek alkalmazását a képességek teljes életciklusára kiterjeszti. A stratégia célja nem a technológia „befékezése”, hanem annak biztosítása, hogy a fejlesztés és alkalmazás összhangban legyen az értékekkel, normákkal és nemzetközi joggal; továbbá kezelje a rosszindulatú felhasználásból eredő fenyegetéseket.[67]

A szövetségi műveletekben a policy-k jelentősége a közös bizalom: a partnernemzeteknek érteniük kell, milyen elvek és kontrollok mentén működik egy MI-alapú képesség, különösen, ha az információmegosztás vagy a közös célprioritás függ tőle.

Az USA DoD RAI megközelítésben a felelősség, a jogszerűség és a privacy/civil liberties védelme konkrét governance-mintákban is megjelenik (irányítási szerepkörök, felügyeleti fórumok, T&E, dokumentáció). Ez a logika a NATO-szintű interoperabilitásban is releváns: a koalíciós műveletekben a rendszerek kockázatainak és korlátainak transzparens kommunikálhatósága előfeltétele a közös alkalmazásnak.[21], [67]

Az EU AI Act a magas kockázatú rendszerek esetén részletes követelményeket ír elő (kockázatkezelési rendszer, adatirányítás, technikai dokumentáció, naplózás, átláthatóság és emberi felügyelet, valamint pontosság, robosztusság és kiberbiztonság). A katonai fejlesztésekben a rendelet közvetlen alkalmazhatósága korlátozott lehet, ugyanakkor a követelmények szerkezete erős „audit-logikát” ad: mit kell igazolni és milyen bizonyítékokkal. Ez különösen a beszállítói lánc, a dual-use komponensek és a civil-katonai átmenetek miatt válhat relevánssá.[64]

9.2.1. Policy-keretek közös nevezői és feszültségpontjai

A policy- és standard-keretek közös metszete a következő: (1) kormányzás és felelősség; (2) emberi felügyelet; (3) átláthatóság, nyomonkövethetőség és magyarázhatóság; (4) megbízhatóság, robosztusság és biztonság; (5) rosszindulatú felhasználás és ellentevékenység kezelése. A feszültségpontok jellemzően a kötelező erő (jogszabály vs. ajánlás), a megfelelőség igazolásának módja (conformity assessment, audit), valamint a katonai/nemzetbiztonsági kivételek és a transzparencia szükséges szintje körül jelennek meg.[12], [21], [64], [67]

9.2. táblázat – NATO/EU/USA/NIST/ICRC keretek összehasonlítása (katonai MI)

Keret	Fókusz	Kiemelt követelmények	Katonai relevancia a COPD-ben
DoD RAI (2021)	Etikai elvek és governance a DoD életciklusban	Etikai elvek; felügyelet; T&E; privacy/civil liberties; dokumentáció	Beszerezés és képességfejlesztés; T&E és üzemeltetés; felelősségi rend a törzsnél
NIST AI RMF (2023)	Általános kockázatkezelési modell	Govern–Map–Measure–Manage; metrikák; folyamatos monitorozás	Kockázati jegyzék és mérési rendszer a tervezés–végrehajtás–értékelés láncban
NATO AI Strategy (2021)	Felelős használat elvei szövetségi keretben	Responsible Use elvek; életciklus-szemlélet; rosszindulatú használat kezelése	Interoperabilitás és bizalom; közös elvek a műveleti alkalmazásban
EU AI Act (2024)	Kockázatalapú jogi szabályozás	High-risk követelmények: dokumentáció, naplózás, human oversight, cybersecurity	„Best practice” megfelelési logika; dual-use és civil-katonai átmenetek
ICRC LAWS diskurzus	Human control és IHL megfelelés autonómia mellett	Kritikus funkciók; előreláthatóság; korlátozás és szabályozás	Célkijelölés/tűzengedély közeli MI: emberi kontroll és jogi kockázatok

9.2.2. A „kritikus funkciók” és a felelősség problémája

A nemzetközi viták egyik visszatérő eleme, hogy hol húzódik a határ a döntéstámogató és a döntést „kiváltó” automatizáció között. A kritikus funkciókhoz (cél kiválasztása és támadás) közelítő rendszereknél az ICRC dokumentumai az előreláthatóság, a korlátozhatóság és a megfelelő emberi kontroll megtartását tartják központi kérdésnek.[69] A katonai tervezésben ez lefordítható a „mely döntési pontokhoz engedjük közel az MI-t” kérdésre: például a cél-lista előkészítése (prioritásjavaslat) lehet MI-vel támogatott, de a jogi megfelelési és arányossági vizsgálatot, valamint a tűzengedélyt emberi döntéshez kell kötni.

A felelősség kérdése a modern, komplex rendszerekben több szereplő között oszlik meg: fejlesztő, beszállító, integrátor, üzemeltető szervezet, parancsnok és operátor. A DoD RAI és a NIST AI RMF közös üzenete, hogy ezeket a felelősségi viszonyokat nem lehet „hallgatólagosan” kezelni; a governance struktúrát és a döntési jogosultságokat dokumentáltan

kell rögzíteni, és a rendszer működését úgy kell kialakítani, hogy a felelősségi lánc utólag is rekonstruálható legyen (auditability).[21], [12]

A szövetségi műveletekben külön kihívás, hogy a felelősségi lánc több nemzeti rendszer és policy kereszteződésében működik. A NATO AI Strategy elvei ezért olyan minimumot jelölnek ki, amely a partnernemzetek között alapvető interoperabilitási és bizalmi keretet ad.[67] A gyakorlatban ez a közös terminológia, a közös T&E elvárások, valamint a kockázati információk megosztása (pl. model limitations, known failure modes) formájában jelenhet meg.

9.3. Adatvédelem, torzítás, explainability, félreinformálás és manipulációs kockázatok

A katonai MI-rendszerek teljesítménye és kockázata döntően az adaton múlik. A modern döntéstámogatás gyakran személyes adatokkal (pl. biometrikus azonosítók, eszköz- és helyadatok, kommunikációs metaadatok), érzékeny OSINT-adatokkal, valamint minősített ISR-információval dolgozik; a generatív MI esetén pedig a felhasználói promptok, a chat-naplók és a rendszer által generált összefoglalók is információhordozók. Ez adatvédelmi és adatbiztonsági szempontból kétirányú kockázat: egyrészt a jogosulatlan hozzáférés és kiszivárgás, másrészt az adatok „másodlagos felhasználása” (purpose creep), amikor az adatokat a kezdeti célon túl, nem átlátható módon használják fel.[21], [12], [66]

A NIST AI RMF és az EU AI Act egyaránt hangsúlyozza az adatirányítást, az adatminőség és a dokumentáció szerepét. A katonai alkalmazásban ez lefordítható olyan kérdésekre, mint: milyen forrásból származik az adat, mennyire reprezentatív az adott OE-re, milyen minőségi problémákat és torzításokat ismerünk, valamint milyen hozzáférési és műveleti korlátozások védik. Kockázatkezelési szempontból a „nem tudjuk, milyen adatból tanult a modell” állapot önmagában kritikus: ellehetetleníti a hibamódok előrejelzését, és gyengíti az auditálhatóságot.[12], [64]

A torzítás (bias) a katonai döntéstámogatásban különösen veszélyes, mert hatása kumulatív: torz észlelés vagy előrejelzés torz erőforrás-elosztáshoz, torz célprioritáshoz, végső soron pedig aránytalan kockázatvállaláshoz vezethet. A bias forrása lehet mintavételi torzítás (csak „látványos” események kerülnek be), mérési torzítás (szenzorok, annotáció), nyelvi torzítás (alacsony erőforrású nyelvek), valamint strukturális torzítás (a múltbeli döntések öröksége). A

felelős MI keretek ezért követelik a bias azonosítását és kezelését, valamint annak dokumentálását, hogy a rendszer hol és hogyan lehet sérülékeny.[21], [12], [64]

A magyarázhatóság (explainability) és nyomonkövethetőség a katonai környezetben kettős célt szolgál: (1) támogatja a törzs és operátorok kalibrált bizalmát és a helyes használatot; (2) biztosítja a felelősségi és jogi elszámoltathatóságot (auditability). A magas kockázatú rendszereknél a pusztán pontossági metrika nem elég: tudni kell, milyen helyzetekben hibázik a modell, milyen bizonytalansággal ad javaslatot, és milyen „tiltott zónákban” nem alkalmazható. Ez a gondolkodásmód a NIST AI RMF mérési és menedzsment funkcióiban, valamint az EU AI Act dokumentációs és human oversight követelményeiben is megjelenik.[12], [64].

A félreinformálás és manipuláció kockázatai a generatív MI korszakában a katonai műveletek információs környezetét is átalakítják. A mesterségesen generált tartalmak (szöveg, kép, hang, videó) alacsony költséggel és nagy sebességgel állíthatók elő, miközben a „hitelesség-heurisztikákat” kihasználva növelhetik a megtévesztés esélyét. A von Sikorski–Hameleers áttekintés kiemeli a jelenség kettős felhasználású jellegét: ugyanazok az eszközök, amelyek támogatják a tartalom-előállítást és ellenőrzést, alkalmazhatók a manipuláció skálázására is.[39]

A katonai döntéstámogatás szempontjából a manipuláció nemcsak a lakossági információs műveleteket érinti. A törzs helyzetértékelése OSINT-, HUMINT- és partnerjelentésekből is építhető, amelyek a generatív MI és bot-hálózatok miatt nagyobb arányban tartalmazhatnak szintetikus vagy hamis elemeket. Emiatt a forráskritika, a többforrású megerősítés, valamint a tartalom-eredet (provenance) kezelése az MI-alapú elemző lánc integráns része. A NIST AI 600-1 külön kockázatként kezeli a szintetikus tartalmakból eredő megtévesztést, és javasolja a digitális transzparencia és jelölés (pl. provenance) alkalmazását.[66]

9.3.1. Adatvédelem és adatbiztonság: minősített és nyílt adatok együttélése

A katonai MI-rendszerek gyakorlati bevezetésének egyik legnehezebb pontja a minősített és nem minősített (nyílt) adatok együttélése. A törzsek gyakran egyszerre támaszkodnak minősített ISR-információra és nyílt forrású jelzésekre, miközben a generatív MI-rendszerek „táplálása” (promptok, kontextus) könnyen adatkeveredéshez vezet. A felelős MI keretek

alapján a minimális követelmény az adat-szegmentáció és a hozzáférés-szabályozás: ki, milyen adatszinten, milyen feladatra használhat MI-eszközt.[21], [12]

A privacy/civil liberties dimenzió a katonai környezetben is releváns, különösen stabilizációs műveletekben, városi környezetben és nagy mennyiségű OSINT feldolgozásakor. A DoD RAI dokumentum explicit módon kiemeli, hogy a felelős alkalmazás a privacy védelmét és a jogszerű adatkezelést is magában foglalja.[21] A gyakorlatban ez adatminimalizálást, célhoz kötöttséget, hozzáférési naplózást, valamint törzs-szintű ellenőrzési pontokat jelenthet, ahol a személyes adat kezelése és az MI-eredmények felhasználása felülvizsgálatra kerül.

9.3.2. Torzítás, hibamódok és magyarázhatóság a törzsmunkában

A torzítás kezelése a katonai döntéstámogatásban nem pusztán fairness-probléma, hanem megbízhatósági és kockázatvállalási kérdés. Ha a modell bizonyos régiókban, nyelvekben vagy társadalmi csoportoknál szisztematikusan gyengébb teljesítményt nyújt, akkor a törzs – akár tudtán kívül – aránytalan kockázatot vállalhat, vagy hibás prioritásokat állíthat fel. A NIST AI RMF a káros torzítások kezelését a „trustworthiness” egyik jellemzőjeként tárgyalja, és a mérésben a reprezentativitás és hibaszint szegmentált értékelését javasolja.[12]

A magyarázhatóság gyakorlati szintje a döntési helyzettől függ. A törzs számára sok esetben nem a matematikai mechanizmus részletes magyarázata a legfontosabb, hanem: (1) mely inputok voltak döntőek; (2) mennyire stabil a javaslat kis változások mellett; (3) milyen korábbi példákhoz hasonlít a helyzet; (4) mennyire biztos a rendszer, és miért. Az EU AI Act és a NIST keretek közös üzenete, hogy a magas kockázatú alkalmazásokban az emberi felügyeletet csak akkor lehet komolyan venni, ha a felhasználó releváns magyarázó információt kap.[12], [64]

9.3.3. Manipuláció és dezinformáció: kockázati indikátorok a műveleti környezetben

A dezinformáció elleni védekezésben a pusztán „tartalom-ellenőrzés” ritkán elegendő: a döntéstámogatásnak a terjedési mintákra, a hálózati viselkedésre és a kontextusra is figyelnie kell. A modern információs térben a dezinformáció iparszerű előállítás, automatizált terjesztése és célzott személyre szabása olyan ellenség-képesség, amely a törzs helyzetértékelését és legitimációját egyszerre támadja.[39]

3. táblázat – Információs manipuláció kockázati indikátorai és mitigációk (példák)

Kockázati jelenség	Tipikus indikátorok	Műveleti hatás	Lehetséges mitigáció
Szintetikus tartalom (deepfake, AI-szöveg)	Forráshiány; metaadat-anomáliák; időzített „kampány”; több csatornán azonos narratíva	Téves helyzetkép; reputációs kár; döntések késleltetése	Többforrású validáció; tartalom-eredet vizsgálat; forráslista és bizalmi szintek
Bot/hamis fiók hálózatok	Koordinált posztolás; rendellenes hálózati struktúra; atipikus aktiválási hullámok	Narratív dominancia; polarizáció; műveleti támogatás gyengítése	Hálózatelemzés; anomáliaészlelés; platformokkal együttműködés
„MI generálta” vád delegitimációra	Hiteles tartalom tagadása „AI” címkével; gyors narratív váltás; bizonyítékok relativizálása	Bizalomerózió; döntési bizonytalanság; kommunikációs válság	Transzparens bizonyítéksomag; kommunikációs protokoll; nyílt források archiválása
Adatforrás mérgezése OSINT-ben	Hirtelen „új” források; túl homogén üzenetek; földrajzi/nyelvi anomáliák	Hamis fenyegetéskép; erőforrás-átcsoportosítás; téves prioritások	Forrás-értékelési szabályok; whitelisting; idősoros konzisztenciateszt

A táblázatban jelzett kontrollok akkor működnek, ha a törzs (a) rendelkezik releváns adatokkal a forrásokról és terjedési mintákról, (b) az MI-elemző láncban a bizonytalanság és a forrás-bizalom „látható” marad, és (c) a kommunikációs funkciók összehangoltan reagálnak a bizalom elleni támadásokra. A dezinformáció kockázata ezért egyszerre technikai, szervezeti és kommunikációs kérdés.[66], [39].

9.4. Adversarial ML, modellmérgezés, ellentevékenységek; védekezési stratégiák

A katonai MI-rendszerek ellenséges környezetben működnek, ezért a kockázatmodellnek alapértelmezésben számolnia kell szándékos támadással. Az adversarial machine learning (AML) gyűjtőfogalomként írja le azokat a módszereket, amelyek a modell hibáit szándékosan váltják ki (evasion), a tanítási folyamatot mérgezik (poisoning), a modellt rekonstruálják vagy „ellopják” (extraction), illetve a bemeneti csatornán (például prompt) keresztül nem kívánt viselkedést idéznek elő. Brundage és munkatársai áttekintése arra mutat rá, hogy a

„megbízhatóság” és „biztonság” nem választható el az adversarial környezet kezelésétől: a támadó képes a kockázati profil megváltoztatására még azonos algoritmus mellett is.[71]

A klasszikus adversarial példák – például képosztályozók félrevezetése alig észrevehető perturbációval – azt jelzik, hogy a magas pontosságú modellek is sérülékenyek célzott inputokkal szemben. Goodfellow és szerzőtársai a jelenséget a modellek lokális linearitásával és a nagy dimenziószámú terekben fellépő érzékenységgel magyarázzák. Ez a katonai számítógépes látás (CV) alkalmazásokban – célfelismerés, járműazonosítás, szenzorfüzió – közvetlen műveleti kockázatot jelenthet.[72]

A generatív nyelvi modellek és asszisztensek esetén az ellentevékenységek gyakran prompt-injekció, kontextus-mérgezés vagy adat-kiszivárgás formájában jelennek meg: az ellenség olyan bemenetet juttat a rendszerbe, amely a szabályokat megkerülve érzékeny információt szivároztat, hibás ajánlást generál, vagy a felhasználót manipulálja. A NIST AI 600-1 kifejezetten tárgyalja a prompt-alapú támadásokat, a nem kívánt tartalomgenerálást és a modellkiszivárgás kockázatát, ami a katonai környezetben a titokvédelem és az OPSEC szempontjából kritikus.[66]

Modellmérgezés (poisoning) nemcsak a tréningadatokba való szándékos beavatkozást jelenti, hanem az ellátási lánc kockázatát is: kompromittált adatkészletek, sérülékeny könyvtárfüggőségek, nem ellenőrzött előtanított modellek, valamint a konfigurációk és súlyfájlok integritásának hiánya. A NIST AI RMF kockázatkezelési logikája szerint ezek a fenyegetések csak akkor kezelhetők, ha a szervezet képes feltérképezni (Map) az MI-rendszer teljes ökoszisztémáját és mérni (Measure) a kontrollok hatékonyságát.[12]

A védekezési stratégiák ezért több rétegben értelmezendők: (1) adatbiztonság és integritás (input- és tréningadat); (2) modell-robosztusság (például adversarial teszt és hardening); (3) futásidejű védelem (anomáliaészlelés, sandbox, policy engine); (4) hozzáférés- és jogosultságkezelés (least privilege, zero trust); (5) szervezeti red teaming és folyamatos monitorozás. A DoD RAI és a NIST keretek közös üzenete, hogy a védelem nem telepítéskor kezdődik: a követelményekben és a T&E-ben kell rögzíteni a minimális biztonsági szintet.[21], [12], [71]

9.4.1. AML fenyegetési térkép: támadási felületek és védelmi kontrollok

A katonai MI rendszerek fenyegetési modellje célszerűen a támadási felületek szerint strukturálható. A felületekhez hozzárendelt kontrollok lehetővé teszik a „defense in depth” megvalósítását, és segítenek elkerülni, hogy a kiberbiztonság kizárólag hálózati vagy infrastruktúra-kérdéssé redukálódjon.[12], [71]

9.4. táblázat – Adversarial ML fenyegetések és védekezési stratégiák (összefoglaló)

Támadási felület	Támadási típus (példa)	Várható hatás	Védelmi stratégia (példák)
Bemeneti adatok (szenzor, OSINT)	Evasion / adversarial input	Hamis osztályozás; téves riasztás; téves célazonosítás	Input-szűrés; többforrású megerősítés; adversarial teszt; redundancia
Tréning és finomhangolás	Data poisoning; backdoor	Rejtett hibák; „trigger” aktiválás; torz döntéstámogatás	Adat-provenance; integritás-ellenőrzés; kontrollált finomhangolás; auditálható pipeline
Modell és súlyfájlok	Model tampering; supply chain	Kimenetek manipulálása; paraméterek módosítása	Digitális aláírás/ellenőrzés; verziókezelés; biztonságos tárolás; hozzáférés-korlátok
Interfész (chat, API)	Prompt injection; jailbreak; data leakage	Érzékeny adatok kiszivárgása; hibás ajánlások; manipuláció	Policy engine; input/output szűrés; sandbox; titokvédelmi szabályok
Üzemeltetés	Drift kihasználása; monitoring megkerülése	Fokozatos teljesítményromlás; „csendes” hibák	Folyamatos monitorozás; drift detektálás; küszöbök; rollback

9.4.2. Kiberbiztonság és MI: integrált incidenskezelés és tanuló védelem

A hagyományos kiberbiztonsági incidenskezelés és az MI-üzemeltetés összekapcsolása alapvető feltétel a katonai alkalmazásokban. Ha a rendszer teljesítménye romlik, vagy rendellenes kimeneteket ad, az lehet modell-drift, lehet adathiba, de lehet ellenséges manipuláció is. A NIST AI RMF ezért a Manage funkció alatt hangsúlyozza a folyamatos

monitorozást, a riasztási küszöbök kijelölését és a korrigáló intézkedések (mitigation) „előkészített” jellegét.[12] A gyakorlatban ez azt jelenti, hogy a törzs számára előre definiált SOP szükséges arra, mikor kell a rendszert korlátozott üzemmódra kapcsolni, mikor kell visszatérni manuális eljárásokhoz, és milyen bizonyítékokat kell megőrizni a vizsgálathoz.

A generatív MI rendszerekben az incidens fogalma kiterjedhet a titokvédelmi és információbiztonsági eseményekre is: például nem kívánt adatmegjelenítés, túl sok „belső” kontextus visszaadása, vagy olyan válaszok, amelyek a felhasználót félrevezetik. A NIST AI 600-1 profil a generatív rendszerek esetében kifejezetten javasolja a kimenet-szűrést, a felhasználói interakciók naplózását, valamint a „model governance” beépítését az üzemeltetésbe.[66]

Az integrált védelem szervezeti feltétele, hogy a kiberbiztonsági szakemberek, az adat/MI csapat és a műveleti felhasználók közös fenyegetési modellt használjanak. Brundage és munkatársai rámutatnak, hogy a támadási felületek heterogének: az ellenség nem csak a hálózatot, hanem a modellt, az adatokat és a felhasználói interfészt is támadhatja.[71] Ennek megfelelően a „red teaming” és a rendszeres adversarial teszt a katonai T&E (tesztelés-értékelés) része kell legyen, nem pedig ad hoc kutatási aktivitás.

9.5. Minőségbiztosítás, tesztelés-értékelés (T&E), megfelelés és audit

A felelős MI a gyakorlatban akkor válik működőképessé, ha a minőségbiztosítás (QA) és a tesztelés-értékelés (T&E) nem utólagos formalitás, hanem a fejlesztési és beszerzési döntésekbe épül. A DoD RAI dokumentum a T&E szerepét központi elemként kezeli: a rendszernek igazolhatóan megbízhatónak, nyomonkövethetőnek és irányíthatónak kell lennie, különösen a műveleti alkalmazásokban.[21]

A NIST AI RMF a Measure és Manage funkciókban részletezi, hogyan kell metrikákat és értékelési eljárásokat hozzárendelni a kockázatokhoz. Ez katonai környezetben azt jelenti, hogy nem elég „benchmark” pontosságot mérni; a tesztelésnek a műveleti környezet realitásait (zaj, ellentevékenység, adatdrift, hiányos megfigyelés, adversarial inputok) is reprezentálnia kell. Különösen fontos a hibamódok dokumentálása és a „known limitations” kommunikálása a törzs felé.[12], [71]

Az EU AI Act logikája szerint a magas kockázatú rendszerek megfelelőségének bizonyítása dokumentációval és eljárásrenddel történik: technikai dokumentáció, naplózás, kockázatkezelési rendszer, valamint emberi felügyelet és kiberbiztonsági garanciák. A katonai rendszerekben – még ha nem is közvetlenül jogszabályi megfelelés okán – ez a logika azért hasznos, mert egy koalíciós vagy nemzeti audit, illetve egy művelet utáni vizsgálat esetén a „bizonyítékcsoomag” megléte határozza meg, mennyire rekonstruálható a döntési folyamat.[64]

A generatív MI alkalmazásoknál a minőségbiztosítás kiterjed a tartalombiztonságra és a nem szándékolt kimenetek kezelésére. Az AI 600-1 profil többek között javasolja a tesztadatok és prompt-készletek kialakítását tipikus támadási és hibamintákra (prompt injection, hallucination/confabulation, data leakage), valamint a működési korlátok meghatározását (mire nem használható a rendszer). A katonai alkalmazásban ez különösen releváns a titokvédelem és az OPSEC miatt.[66]

Az auditálhatóság gyakorlati feltétele, hogy a rendszer minden kritikus komponense (adat, modell, konfiguráció, döntési log) verziózott és visszakereshető legyen. Az audit nem feltétlenül külső hatósági ellenőrzést jelent; gyakran belső ellenőrzés, utólagos vizsgálat (after-action review) és tanulási folyamat. A RAI és AI RMF keretek e téren közös nevezőre jutnak: a kockázatok kezelésének bizonyítékokon kell alapulnia, és a bizonyítékot a működés során is folyamatosan gyűjteni kell.[21], [12]

9.5.1. T&E ellenőrzőlista katonai döntéstámogató MI-re

A tesztelés-értékelés gyakorlati hasznát egy olyan ellenőrzőlista adja, amely a rendszer életciklusát követi. A lista célja nem a kreatív mérnöki megoldások korlátozása, hanem annak biztosítása, hogy a törzs és a fejlesztő szervezet ugyanazt tekintse minimális bizonyítási szintnek a telepítés előtt és alatt. A lista összehangolható a DoD RAI életciklus-szemléletével, a NIST AI RMF funkcióival és az EU AI Act dokumentációs logikájával.[21], [12], [64]

9.5. táblázat – MI T&E és megfelelőségi ellenőrzőlista (példa)

Életciklus-fázis	Mit kell tesztelni/értékelni?	Metrika / példa mérés	Bizonyíték
Követelmények	Felhasználási korlátok, tiltott alkalmazások, döntési felelősség	Use-case és misuse- case elemzés; kockázati tolerancia	Követelményjegyzék; risk register

Adat	Forrás, reprezentativitás, minőség, torzítás, hozzáférés	Hiányosságok; bias riport; integritás-ellenőrzés	Adatleltár; adatlappal bővített dataset
Modell	Teljesítmény + bizonytalanság + robosztusság	Stresszteszt; adversarial teszt; kalibráció	T&E riport; model card
Integráció	Interfész, naplózás, jogosultság, fail-safe	Log lefedettség; hozzáférés tesztek; fallback gyakorlat	Integrációs tesztjegyzőkönyv; konfigurációs jegyzék
Üzemeltetés	Drift, incidenskezelés, frissítések, audit	Drift-riasztások; SLA; red teaming eredmények	Monitoring riportok; incidensnapló; auditjelentés

9.5.2. A megfelelés bizonyítása: dokumentáció, naplózás, „evidence by design”

A megfelelés és auditálhatóság megtervezése („evidence by design”) azt jelenti, hogy a rendszer már a tervezéskor úgy készül, hogy bizonyítékokat termeljen: strukturált naplózás, verziókövetés, konfiguráció-kezelés, és olyan döntési naplók, amelyek összekapcsolják a kimeneteket az inputokkal és a felhasználói döntésekkel. Ez a logika közvetlenül megjelenik az EU AI Act dokumentációs és naplózási követelményeiben, és összhangban van a DoD RAI és a NIST AI RMF governance fókuszával is.[21], [12], [64]

A katonai környezetben az „evidence” gyakran minősített, ezért a bizonyíték-kezelésnek figyelembe kell vennie a titokvédelmi szabályokat és a hozzáférési jogosultságokat. Ugyanakkor a bizonyíték hiánya a művelet utáni értékelést és a felelősségi vizsgálatot teszi nehezzé. A gyakorlatban ezért célszerű két szinten kezelni a bizonyítékot: (1) operatív, minősített bizonyítékcsoomag a konkrét műveleti döntésekhez; (2) aggregált, anonimizált vagy „metaszintű” bizonyítékcsoomag a modell teljesítményének és kockázatainak követéséhez. Ez összhangban áll a NIST AI RMF megközelítésével, amely a kockázatkezelést szervezeti tanulási folyamatként is értelmezi.[12]

9.6. Összefoglalás

A fejezet bemutatta, hogy az MI katonai alkalmazása a COPD teljes spektrumán a teljesítménynövekedés mellett új kockázatokat hoz: jogi és etikai (IHL, felelősség, emberi kontroll), információs (dezinformáció, manipuláció), technikai (roboztusság, drift), valamint kiberbiztonsági (adversarial ML, ellátási lánc) dimenziókban. A felelős MI nem választható el

a műveleti hatékonyságtól: a bizalom, az elszámoltathatóság és az interoperabilitás feltétele, hogy a rendszerek kockázatai kezelhetők és bizonyíthatók legyenek.[12], [21], [64], [67]

A NIST AI RMF és a generatív profil (AI 600-1) módszertani gerincet ad a kockázatok feltérképezésére, mérésére és menedzselésére; a DoD RAI és a NATO AI Strategy pedig a katonai governance és értékalapú alkalmazás irányát jelöli ki.[12], [21], [67], [66] A nemzetközi diskurzus (ICRC, CCW GGE LAWS) rámutat arra, hogy a human control és a kritikus funkciók kezelése a legitim alkalmazás egyik központi kérdése marad.[69], [70]

A gyakorlati ajánlás a COPD szemszögéből: a kockázatkezelést és a megfelelőségi bizonyítékokat már a képesség követelményeibe (COA/CONOPS/OPLAN) be kell ágyazni, majd a végrehajtás során folyamatosan monitorozni kell. Az MI-t csak olyan döntési pontokhoz érdemes közel engedni, ahol a felügyelet, a naplózás, a robosztussági tesztek és az incidenskezelés arányos módon biztosítható. Ezzel érhető el, hogy az MI a parancsnoki szándékot erősítse, ne pedig új, nehezen kontrollálható kockázatokat termeljen.[21], [12], [66]

10. ÖSSZEGZETT KÖVETKEZTETÉSEK

Jelen disszertáció központi kérdése az volt, hogy a mesterséges intelligencia (MI) milyen módon, milyen tervezési pontokon és milyen szervezeti-technológiai feltételek mellett tudja érdemben támogatni a NATO Comprehensive Operational Planning Directive (COPD) művelettervezési folyamatát. A vizsgálat fókuszában a COPD teljes életciklusa állt: (i) a műveleti környezet megértésétől és a stratégiai helyzetértékeléstől, (ii) a katonai válaszopciók, a koncepció és az OPLAN kidolgozásán át, (iii) a végrehajtás és az átmenet adaptív irányításáig, valamint (iv) a kockázatok, az etikai-jogi megfelelés és a kiberbiztonság kezeléséig.

A kutatás alapját a COPD v3.1 folyamatszemplélete adta: a tervezési termékek (pl. helyzetértékelések, COA/CONOPS/OPLAN dokumentumok, döntési briefek és végrehajtási visszacsatolások) olyan információs objektumokként értelmezhetők, amelyek előállítása, karbantartása és ellenőrzése mind emberi, mind gépi (MI-alapú) erőforrásokat igényel. Ennek megfelelően a disszertáció nem egyetlen algoritmus vagy alkalmazás bemutatására törekedett, hanem egy rendszerszintű integrációs szemlélet felvázolására: hogyan lehet az MI-képességeket a törzsmunkába úgy beépíteni, hogy közben teljesüljenek a megbízhatóság, az átláthatóság, az auditálhatóság és az információbiztonság követelményei.[73], [12]

A stratégiai és szövetségi környezet gyors ütemű digitalizációja, a többdoménű műveletek adat-intenzív jellege, valamint a hálózatba kapcsolt szenzorok és döntéshozók igényei olyan képességfejlesztési irányokat jelölnek ki, amelyekkel a művelettervezés MI-támogatása szorosan összefügg. A Joint All-Domain Command and Control (JADC2) gondolatkör, valamint a NATO digitális transzformációs céljai közös metszetben kezelik a gyors adatáramlást, a közös helyzetkép és a döntéstámogatás követelményeit.[58], [74]–[76]

10.1. Az elvégzett vizsgálat tömör leírása és fejezetenként összegzett következtetésem

A disszertáció a művelettervezést – a klasszikus törzseljárások tiszteletben tartása mellett – részben információfeldolgozó és döntéstámogató rendszerként modellezte. Ez a megközelítés lehetővé tette, hogy a COPD egyes fázisaihoz (1–6) konkrét MI-funkciókat rendeljünk: adatgyűjtést és -tisztítást, össz-adatforrású adatfúziót, elemzést, predikciót, COA generálást és összevetést, erőforrás-optimalizációt, szimulációt, valós idejű monitorozást, valamint utólagos (after action) értékelést.

A fejezetek végkövetkeztetései a következő logika szerint összegezhetők: (i) a doktrínák és elméleti keretek kijelölték a művelettervezés „kötelező” döntési és dokumentációs pontjait, (ii) az MI ott képes a legnagyobb hozzáadott értéket adni, ahol az információ mennyisége nagy, az idő kritikus, és az emberi feldolgozó kapacitás szűk keresztmetszet, (iii) az MI csak rendszerszinten – adatstratégiával, kockázatkezeléssel és minőségbiztosítással összekapcsolva – vezet fenntartható képességnövekedéshez.

10.1.1. Elméleti keret: a tervezéstámogató MI minimum-követelményei

A 2. fejezetben kialakított elméleti keret alapján azt a következtetést vontam le, hogy a katonai művelettervezést támogató MI-alkalmazásoknál nem elegendő a „jó pontosság” vagy a „gyors futásidő” elérése. A tervezéstámogató MI-nek a döntéshozói és törzsmunkai környezetben minimumként biztosítania kell: (i) az adatok és a modellek származásának (provenance) és bizonytalanságának követhetőségét, (ii) a magyarázhatóság olyan szintjét, amely támogatja a szakértői ellenőrzést, (iii) az emberi kontrollt és a felelősségi körök tiszta kijelölését a humán–gép csapatban, (iv) a torzítások (bias) és a manipulációk (deception, adversarial ML) elleni ellenállóképességet, valamint (v) a kiberbiztonsági és minősített adatkezelési előírások betartását.

Az elméleti keretben az MI-t nem „autonóm döntéshozóként”, hanem „tervezési erőforrásként” értelmeztem: olyan eszközként, amely a szakértői ítéletalkotást gyorsítja és kiterjeszti. Ennek kritikus következménye, hogy az MI-támogatás értékelésénél a rendszer egészének teljesítményét kell mérni (pl. döntési ciklusidő, produktumok minősége, hibák csökkenése), nem pusztán az egyes modellek technikai metrikáit.

A 2. fejezet megalapozta azt a későbbi állításomat is, hogy a felelős MI (Responsible AI) és a kockázatkezelés nem „külön fejezet”, hanem a teljes tervezési láncba beágyazott követelmény. Ezért a későbbi fejezetekben a kockázatkezelési szempontokat a tervezési termékekhez és a döntési pontokhoz rendeltem (pl. COA-választás, végrehajtási adaptáció).

10.1.2. A katonai művelettervezés fejlődése: miért időszerű az MI-integráció?

A 3. fejezet történeti áttekintése alapján arra jutottam, hogy a művelettervezés evolúcióját mindig egyszerre formálta: (i) a hadműveleti gondolkodás (operational art) fejlődése, (ii) a szervezeti tanulás (törzsfolyamatok, standardizált termékek), valamint (iii) az információs és

számítási képességek bővülése. A modern, többdoménű műveletekben az információ mennyisége, heterogenitása és a helyzet dinamikája olyan mértéket ért el, amelyben a klasszikus törzseljárások önmagukban már nem garantálnak megfelelő reakciósebességet.

A fejezetből levont részkövetkeztetésem, hogy az MI integrációja nem „forradalmi szakítás”, hanem a tervezés hosszú távú trendjeinek következő lépcsője: a tervezők munkáját támogató eszközök (pl. adatbázisok, térinformatikai rendszerek, szimulációk) után az MI a mintázatfelismerést, a predikciót és a többváltozós optimalizációt emeli be a mindennapi törzsmunkába. Ugyanakkor a történelem azt is mutatja, hogy minden új eszköz csak akkor hozza a várt hatást, ha a doktrína, a képzés és a szervezeti kultúra is adaptálódik.

Ebből következően a disszertáció nemcsak technológiai, hanem szervezeti ajánlásokat is megfogalmaz (képességfejlesztési útiterv, interoperabilitás, governance), amelyek nélkül a „pilot jellegű” MI-megoldások nem skálázhatók műveleti szintre.

10.1.3. NATO és EU tervezési doktrínák: közös pontok és követelmények az MI számára

A 4. fejezetben a NATO és EU (CSDP) művelettervezési megközelítéseit összevetve arra jutottam, hogy a tervezési rendszerek közös magja a többdimenziós helyzetértékelés, a koalíciós koordináció és a dokumentált döntés-előkészítés. Ez a közös mag egyben az MI-integráció „kényszerpályáját” is kijelöli: az MI-támogatásnak illeszkednie kell a szabványosított tervezési termékekhez, és különösen a többnemzeti törzsmunkához szükséges átláthatósági és elszámoltathatósági követelményekhez.

A fejezet részkövetkeztetése, hogy a koalíciós művelettervezésben az interoperabilitás nem csupán technikai interfész-kérdés. Szemantikai és eljárásrendi interoperabilitás is kell: azaz közös fogalmi készlet (ontológiák, adatmodellek), közös minősítési és adatkezelési eljárások, valamint a tervezési „terméklogika” közös értelmezése (pl. mit jelent egy COA összehasonlítási mátrix, milyen bizonytalanságot hordoz a helyzetértékelés, hogyan kezeljük az eltérő nemzeti korlátozásokat).

A 4. fejezetből következően az MI-alkalmazások tervezésénél már a kezdetektől be kell építeni a koalíciós működésre jellemző constraints-eket: többnyelvűség, heterogén adatforrások, különböző minősítési szintek és a Mission Partner Environment jellegű együttműködések.

10.1.4. COPD v3.1 elemzés: MI-beavatkozási pontok a tervezési termékek mentén

Az 5. fejezet részletes COPD-elemzése alapján azt a következtetést vontam le, hogy a művelettervezés MI-támogatása akkor tud fenntarthatóan működni, ha a fejlesztés a COPD termék- és döntéspont-logikájára épül. A COPD fázisai között nemcsak időrendi, hanem információs-feldolgozási függések is vannak: az 1–2. fázis helyzetértékelései meghatározzák a 3–4. fázis COA-generálásának kereteit, míg az 5–6. fázis végrehajtási adatai visszacsatolják a feltételezéseket és a kockázatokat.

Az 5. fejezet fő megállapítása, hogy a COPD strukturált termékei (pl. elementáris feladatlebontások, hatásláncok, erőforrás- és időparaméterek, kockázati jegyzékek) alkalmasak arra, hogy egy MI-rendszer számára „géppel olvasható” és auditálható reprezentációként szolgáljanak. Ez azonban adatmodell-fegyelmet és „digitális tervezési fonalat” (digital thread) igényel: a követelmények, feltételezések, döntések és végrehajtási visszacsatolások nyomon követhetőségét.

A fejezet további következtetése, hogy a COPD-ben a legnagyobb MI-hozzáadott érték általában nem a végső döntés automatizálása, hanem a döntés-előkészítő alternatívák gyors és transzparens előállítása (pl. COA-k jelöltlistája, érzékenységvizsgálat, „what-if” futtatások) és a bizonytalanság menedzselése.

10.1.5. COPD 1–2. fázis: MI a műveleti környezet megértésében és a stratégiai helyzetértékelésben

A 6. fejezetben kimutattam, hogy a COPD 1–2. fázisaiban (Initial Situational Awareness, Strategic Assessment) a kritikus kihívás az információs túlterhelés és a források heterogenitása. A PMESII-PT/DIME és kapcsolódó keretek (ASCOPE/ICR2), a JIPOE logika, valamint a össz-adatforrású felderítés (HUMINT, SIGINT, GEOINT/IMINT, OSINT, CYBINT) olyan adattérképet alkotnak, amelyben az MI elsődleges szerepe a strukturálás, a mintázatfelismerés és a relevancia-szűrés.

A fejezet részkövetkeztetése, hogy a helyzetértékelés MI-támogatása akkor éri el a kívánt hatást, ha a következő négy képesség együtt jelenik meg: (i) adatminőség és forráskritika (beleértve a dezinformáció és deception kezelését), (ii) össz-adatforrású adatfúzió és bizonytalanságkezelés, (iii) prediktív elemzések a fenyegetésértékelésben, valamint (iv) a törzs

számára használható, bizonyítékalapú termékek előállítása (briefek, térképi rétegek, indikátorok és figyelmeztetések).

A 6. fejezet legfontosabb következtetése számomra, hogy a stratégiai helyzetértékelésben az MI „minőségbiztosítási” funkciója legalább olyan fontos, mint a sebesség: a modell által javasolt állításoknál kötelező a források és a bizonytalanság jelölése, különben az MI a hibákat is nagy sebességgel teríti szét a döntési láncban.

10.1.6. COPD 3–4. fázis: MI a COA/MRO, CONOPS és OPLAN kidolgozásában

A 7. fejezet alapján azt a következtetést vontam le, hogy a COPD 3–4. fázisaiban (Military Response Options, COA, CONOPS, OPLAN) az MI legnagyobb potenciálja a tervezési alternatívák térbeli-időbeli és erőforrásbeli kombinatorikus robbanásának kezelésében áll. A COA-generálás, a szimuláció és a többcélú optimalizáció egymást erősítő képességek: a generálás sok jelöltet állít elő, a szimuláció a bizonytalan kimeneteket becsli, az optimalizáció pedig a korlátok és célok mentén „rendez” a jelöltek között.

A fejezet részkövetkeztetése, hogy a wargaming eszközökkel való integráció kulcsfontosságú: a COA-k értékelése nem tekinthető pusztán matematikai optimalizációs feladatnak, mert a műveleti környezet adaptív és ellenségvezérelt. A sztochasztikus, agent-based és játékelméleti elemeket is tartalmazó szimulációs hurkok képesek a COA-k robusztusságát és érzékenységét láthatóvá tenni.

A fejezetben prototípus-szemléletben bemutattam, hogy a COA-generátor, az optimalizáló modul és a szimulációs hurok „tervezési emlékezetet” is képezhet: egy olyan adatbázist, amely később gyorsítja az újratervezést és a hasonló helyzetek analógiás kezelését. Ezzel összhangban a 7. fejezet legfontosabb korlátként azt azonosította, hogy az MI minősége erősen függ a bemeneti adatok és a szabályrendszerek minőségétől, ezért a validáció és a red teaming alapkövetelmény.

10.1.7. COPD 5–6. fázis: MI a végrehajtásban, az átmenetben és az adaptív újratervezésben

A 8. fejezet megállapításai alapján azt a következtetést vontam le, hogy a végrehajtási fázisban az MI-támogatás értéke a gyors visszacsatolási hurkokban jelenik meg. A tervvalidáció, az adaptív újratervezés és a valós idejű monitorozás olyan területek, ahol a változó környezet

(ellenség, civil környezet, időjárás, logisztika, információs tér) folyamatosan felülírja a korábbi feltételezéseket.

A fejezet részkövetkeztetése, hogy a végrehajtás MI-támogatásához elengedhetetlen egy jól definiált „műveleti mérőrendszer”: milyen indikátorok alapján tekintünk egy hatást elértnek, milyen esemény vált ki újratervezést, és milyen szabályok szerint frissülnek a kockázati nyilvántartások. Az MI itt elsősorban (i) anomáliaészlelésben, (ii) trend- és előrejelző analitikában, (iii) erőforrás-újratervezésben, valamint (iv) művelet utáni elemzésben és tudásmenedzsmentben tud a törzs munkájához hozzájárulni.

A fejezet kiemelte azt is, hogy a végrehajtási adatok minősége és védelme kritikus: az interoperabilitás, a minősítési szintek kezelése, a szenzoradatok hitelessége és a kiberbiztonság nélkül a valós idejű MI-alapú visszacsatolás megbízhatatlanná válik.

10.1.8. Kockázatok, etika-jog és kiberbiztonság: az MI-támogatás fenntarthatósági feltételei

A 9. fejezetben arra a következtetésre jutottam, hogy a katonai MI alkalmazásának legnagyobb stratégiai kockázata nem kizárólag technikai (pl. pontatlan modell), hanem szervezeti és jogi jellegű is (pl. felelősségi körök homálya, nem megfelelő T&E, adatvédelmi és kiberbiztonsági hiányosságok). A felelős MI elvei (emberi kontroll, megbízhatóság, átláthatóság, kormányzás) csak akkor érvényesíthetők, ha a tervezési lánc minden pontján megjelennek: a képzésben, a folyamatokban, a rendszer-architektúrában és a beszerzésben.

A fejezet részkövetkeztetése, hogy a kockázatok között kiemelt helyen szerepelnek az ellentevékenységek: adversarial ML támadások, modellsérzés, adatmanipuláció, és az információs műveletekhez kapcsolódó félreinformálási kampányok. Ezért az MI-támogatásnak tartalmaznia kell (i) támadásdetektálási és -elhárítási mechanizmusokat, (ii) robusztus tesztelési-értékelési (T&E) protokollokat, (iii) auditálhatóságot és megfelelőségi nyomvonalat, valamint (iv) a minősített és nyílt adatok biztonságos együttélését biztosító architektúrát.

A 9. fejezet összegző megállapítása szerint a NIST AI RMF és a kapcsolódó governance megközelítések alkalmasak arra, hogy a katonai tervezéstámogató MI-t „mérhető kockázatok” mentén vezessük be, és a megfelelőséget ne utólagos adminisztrációként, hanem a rendszer életciklusának részeként kezeljük.[12]

10.1.9. Integrált következtetés: az MI-támogatás értéke és korlátai a COPD egészében

A fejezetenkénti részkövetkeztetések alapján összességében azt állítom, hogy az MI a COPD teljes folyamatában elsősorban a döntés-előkészítés sebességét és minőségét képes javítani, de csak akkor, ha (i) adat- és modellkormányzással, (ii) kiberbiztonsággal és ellentevékenységgel szembeni védelemmel, valamint (iii) standardizált, auditálható tervezési termékekkel együtt kerül bevezetésre. A „tervező MI” hozzáadott értéke a következőkben ragadható meg:

- a helyzetkép gyorsabb és bizonyíték-alapú kialakítása (1–2. fázis),
- több és robusztusabb alternatíva előállítása és összevetése (3–4. fázis),
- adaptív végrehajtás és gyors visszacsatolás (5–6. fázis),
- intézményesített tanulás és tudásmenedzsment (AAR, lessons learned).

Ugyanakkor korlátként azonosítottam, hogy az MI nem képes önmagában kezelni a stratégiai célok és politikai korlátok ellentmondásait, és nem helyettesíti a parancsnoki felelősséget. A fenntartható alkalmazás feltétele az, hogy az MI transzparensen mutassa meg: milyen adatokból, milyen feltételezésekkel és milyen bizonytalansággal jutott egy javaslatához. Ennek hiányában a tervezés legitimációs és megfelelőségi kockázata nő.

Táblázat 10-1. Fejezetenkénti részkövetkeztetések összefoglalása

Fejezet	Fókusz / kérdés	Fő módszerek / keretek	Kulcskövetkeztetés
2.	Elméleti keret és követelmények	RAI, kockázat, tradecraft, ellenállóképesség	Az MI csak emberi kontrollal, bizonytalanság- és forráskezeléssel alkalmazható megbízhatóan.
3.	Tervezés fejlődése	Történeti elemzés, operational art	Az MI a tervezéstámogató eszközök evolúciójának következő lépcsője, de doktrína-képzés nélkül nem skálázható.
4.	NATO/EU doktrínák	Összehasonlító doktrínaelemzés	A koalíciós tervezés átláthatósági és interoperabilitási igényei

			meghatározzák az MI-rendszer követelményeit.
5.	COPD v3.1	Folyamat- és terméke-elemzés	A COPD termékei géppel olvasható adatmodellekre képezhetők; a digital thread alapkövetelmény.
6.	COPD 1–2	PMESII/JIPOE, össz-adatforrású ISR, adatfúzió	A legnagyobb érték a strukturálásban, relevancia-szűrésben, bizonyíték- és bizonytalanságkezelésben van.
7.	COPD 3–4	COA generálás, optimalizáció, szimuláció, wargaming	Az MI a kombinatorikus robbanást kezeli; a robusztusság vizsgálata és red teaming nélkül kockázatos.
8.	COPD 5–6	Monitorozás, adaptív újratervezés, AAR	A valós idejű hurkok gyorsítják a döntési ciklust, de adatbiztonság és interoperabilitás kritikus.
9.	Kockázat és megfelelés	RAI, AI RMF, AML, T&E	A fenntarthatóság feltétele a governance, a kiberbiztonság és a T&E beágyazása a teljes életciklusba.

10.2. Új tudományos eredményeim

A disszertáció keretében a NATO COPD művelettervezési folyamatának MI-támogatását nem elszigetelt eszközök halmazaként, hanem rendszerszintű képességként vizsgáltam. Az alábbi pontokban foglalom össze a kutatás új tudományos eredményeit.

1. **Kidolgoztam** a COPD 1–6 fázisaihoz illesztett MI-funkcióterképet, amely az egyes tervezési termékekhez (SSA, COA/CONOPS/OPLAN, végrehajtási visszacsatolások) konkrét adat- és modellkövetelményeket rendel. A térkép a tervezés információs függőségei alapján jelöli ki az MI-beavatkozási pontokat és a humán kontroll helyét. **Igazoltam a H1 hipotézist**, azaz, hogy a COPD dokumentum-központú tervtermékeinek előállításában az MI-alapú szövegfeldolgozás és tudásmenedzsment

(pl. automatikus kivonatolás, entitás- és relációkinyerés, hivatkozás-konzisztencia ellenőrzés) mérhetően csökkenti a tervezés átfutási idejét, miközben a termékek strukturális megfelelését javítja.

2. **Javaslatot tettem** egy PMESII-PT/DIME-alapú indikátor-katalógus és adatmodell kialakítására, amely a JIPOE logikához és a össz-adatforrású ISR adatfolyamokhoz illeszkedve támogatja a stratégiai helyzetértékelés strukturált, géppel olvasható előállítását és frissítését. **Igazoltam a H2 hipotézist**, azaz, hogy a hírszerzési megalapozást (különösen JIPOE jellegű elemzéseket) támogató adatfűziós és mintázatfelismerő MI-megoldások növelik a helyzetértés mélységét és csökkentik a kognitív torzításokból eredő hibák valószínűségét, feltéve hogy a folyamat humán „human-in-the-loop” ellenőrzéssel működik
3. **Bemutattam** egy COA/MRO generálási és értékelési lánc (generate–simulate–optimize) koncepcionális architektúráját, amelyben a generatív megközelítések, az optimalizáció és a wargaming alapú szimulációk egymást kiegészítve szolgálják a robusztus alternatívák kiválasztását a bizonytalanságok között. **Igazolva a H3 hipotézist**, rámutattam, hogy megfelelő kormányzás és ellenőrzés hiányában az MI-támogatás növeli a megtévesztés, a hibás következtetés és a nem kívánt információszivárgás kockázatát; e kockázatokat csak kockázat-alapú (GOVERN–MAP–MEASURE–MANAGE) keret, auditálhatóság és szigorú adatkezelés mellett lehet elfogadható szintre csökkenteni
4. **Rendszereztem** a végrehajtási fázisban alkalmazható MI-alapú visszacsatolási hurkokat (monitorozás, anomáliaészlelés, adaptív újratervezés, AAR), és ezekhez javasoltam egy mérőszámrendszert, amely a műveleti hatások és indikátorok nyomon követését összekapcsolja a tervezési feltételezések folyamatos validációjával.
5. Integrált kockázat- és megfeleléségi modellt **dolgoztam ki** a tervezéstámogató MI életciklusára, amely a felelős MI elveit (emberi kontroll, átláthatóság, robusztusság) a COPD döntési pontjaihoz rendeli, és a NIST AI RMF logikájával kompatibilis irányítási (governance) és T&E követelményrendszert fogalmaz meg.[12] **Igazoltam a H4 hipotézist**, mivel a COPD-hez illesztett, moduláris, interoperábilis MI-támogató architektúra – amely a NATO digitális transzformációs és adatstratégiai irányokhoz igazodik – növeli a bevezethetőség és elfogadottság esélyét a NATO/EU/MH kontextusban

6. **Rámutattam**, hogy a koalíciós interoperabilitás MI-környezetben nemcsak adatsere, hanem szemantikai és eljárásrendi összehangolás; ennek alapján **javaslatot tettem** a „tervezési digitális fonal” (traceability) és a közös fogalmi készlet (ontológiák, tudásgráfok) alkalmazására a többnemzeti törzsmunkában.
7. Az INCOSE Systems Engineering megközelítésre támaszkodva **felvázoltam** a COPD–PMESII adatforrások–MI technológia rendszerszintű integrációjának SE-keretét: követelménykezelés, architektúra-tervezés, verifikáció-validáció, konfigurációmenedzsment és életciklus-szintű kockázatkezelés összekapcsolását.[77]

10.3. Gyakorlati felhasználási ajánlások

A disszertáció gyakorlati célja az volt, hogy a tervezéstámogató MI-t ne elméleti lehetőségként, hanem bevezethető és működtethető képességként kezelje. A felhasználási ajánlások ennek megfelelően technológiai, szervezeti és policy dimenzióban is megfogalmazódnak.

10.3.1. Az AI technológiák alkalmazása a katonai művelettervezésben

A következő 3–5 évben több olyan feltörekvő MI-technológia várható, amely közvetlenül érintheti a katonai művelettervezést. A legnagyobb hatás azoknál a technológiáknál várható, amelyek egyszerre képesek sok adatforrást kezelni, bizonytalanságot modellezni, és a tervezői munkafolyamatokhoz „beszélgető” vagy „asszisztens” jelleggel kapcsolódni.

1) Nagy nyelvi modellek és multimodális asszisztensek (LLM/MLLM). Ezek akkor adnak értéket, ha a művelettervezési dokumentumok (pl. helyzetértékelések, CONOPS, OPLAN mellékletek) géppel olvasható változatával együtt alkalmazzuk őket. Hasznos felhasználások: gyors kivonatolás és forrásrendezés, követelmény- és feltételezés-ellenőrzés, konzisztencia-vizsgálat a dokumentumok között, valamint többnyelvű koalíciós környezetben a terminológiai egységesítés támogatása.

2) Tudásgráfok, szemantikus keresés és hálózatelemzés. A tervezésben a szereplők, kapcsolatok, hatások és infrastruktúrák hálózatként értelmezhetők. A tudásgráf-alapú

reprezentáció segíti a COA-k „hatásláncainak” és az ellenség rendszerének modellezését, és hidat képez a PMESII-PT indikátorok és a tervezési termékek között.

3) Szimulációval összekapcsolt MI (agent-based, sztochasztikus és hibrid modellek). A COA-értékelésben a szimulációk nem helyettesítik a parancsnoki ítéletet, de képesek a robusztusságot és az érzékenységet kvantifikálni. A wargaming és a szimuláció közötti átjárást erősíti, ha az MI képes a futtatások eredményeit magyarázható módon összefoglalni, és a döntési alternatívák közötti kompromisszumokat Pareto-szemléletben bemutatni.

4) Multi-agent rendszerek és „planning agents”. A speciális célú ügynökök (pl. logisztika, ISR, cyber, információs műveletek) közös keretben, a törzsmunkát tükröző feladatszervezéssel támogatják a tervezést. A kritikus feltétel itt a szerepek, korlátok és felelőségek explicit definiálása, valamint az, hogy az ügynökök által tett javaslatok auditálhatók legyenek.

5) Robusztus optimalizáció és megerősítéses tanulás (RL). A többcélú erőforrás-elosztási feladatokban (pl. ISR kiosztás, készletezés, időzítés) a robusztus és sztochasztikus optimalizáció már rövid távon is alkalmazható. Az RL akkor lehet releváns, ha a környezet dinamikája és az ellenség adaptációja jól modellezhető, illetve ha a tanulási folyamat kontrollált és szimulált térben történik.

E technológiák katonai jelentőségét erősíti az a stratégiai irány, amely az adat- és platformintegrációt (JADC2) és a szövetségi digitális képességek fejlesztését helyezi előtérbe.[74], [75] A Mosaic Warfare és a distributed kill chain koncepciók pedig azt hangsúlyozzák, hogy a gyors, moduláris képesség-összeállítás és a rugalmas „kill web” gondolkodás növeli a műveleti alkalmazkodóképességet.[78], [79]

10.3.2. MH/NATO – eljárásrend, képességfejlesztés, bevezetési útiterv

A Magyar Honvédség és a NATO számára a tervezéstámogató MI bevezetése akkor reális, ha azt képességfejlesztési programként, nem pedig egyedi szoftverbeszerzésként kezeljük. A NATO digitális transzformációs törekvései és az NCIA technológiai irányai olyan keretet adnak, amelyhez nemzeti szinten érdemes illeszkedni.[58], [76]

Javasolt bevezetési logika (szakaszok):

(1) Előkészítés: adat- és folyamatkép (baseline). A tervezési termékek adatmodelljének és a PMESII/JIPOE indikátorok katalógusának kialakítása; adatminőségi szabályok, metaadat-

séma, minősítési és hozzáférési politika rögzítése. E szakasz célja, hogy a törzsmunkában keletkező információk géppel olvasható, később visszakereshető formában álljanak rendelkezésre.

(2) Prototípus és kísérlet: célzott MI-funkciók. Olyan „kis kockázatú” use case-ek kiválasztása, ahol a human-in-the-loop kontroll jól megvalósítható (pl. dokumentum-összegzés, forrásrendezés, anomália-riasztás, COA összevetési táblák előkészítése). A prototípusok értékelését már ekkor T&E keretben kell végezni: mérőszámok, naplózás, red teaming, adatvédelmi és kiberbiztonsági kontrollok.

(3) Integráció: tervezési munkafolyamatok és rendszerek. A sikeres prototípusokat be kell ágyazni a törzs meglévő C2 és tudásmenedzsment eszközeibe, továbbá ki kell alakítani a képességfelelős szervezetet (product owner, data steward, model risk owner, cybersecurity lead). Ez a szakasz kulcs a skálázáshoz: a „pilot” jellegű megoldások ekkor válnak műveleti képességgé.

(4) Skálázás és intézményesítés: többnemzeti interoperabilitás. A NATO környezetben törekedni kell a szabványos interfészekre, közös metaadatokra, és olyan megoldásokra, amelyek a Mission Partner Environment és a több szintű biztonság (MLS) elvárásai mellett is működnek. A hosszú távú cél a tervezési digitális fonal és a közös tanulás: AAR-adatok, lessons learned és tudásgráfok szövetségi szinten.

A fenti szakaszokhoz illeszkedő útitervet a 10-2. táblázat foglalja össze.

Táblázat 10-2. Javasolt bevezetési útiterv a tervezéstámogató MI-képesség kialakításához

Szakasz	Cél	Fő tevékenységek	Kimenetek	Kockázat / mitigáció
1. Előkészítés	Adat- és folyamatbaseline	Adatleltár; metaadat; minősítés; indikátor-katalógus; tervezési termékek adatmodellje	Adatstratégia; governance; „minimum viable data”	Adatminőség hiánya -> data stewardship, minőségi szabályok
2. Prototípus	Célzott use case-ek validálása	LLM asszisztens; forrásrendezés; anomália-riasztás; COA táblák előkészítése;	Mérőszámok; naplók; T&E jelentés; user feedback	Túlzott elvárások -> korlátok és human-in-

		wargaming integráció kísérlete		the-loop szabályok
3. Integráció	Törzsmunkába ágyazás	API-k; jogosultságkezelé s; audit; cyber kontrollok; MLOps	Integrált döntéstámogató modulok; SOP-k; képzés	Shadow IT - > központi üzemeltetés és konfiguráció - menedzsmen t
4. Skálázás	Szövetségi együttműködés	Interoperábilis adatsere; közös ontológia; Mission Partner Environment illesztés	Szövetségi pilot; közös lessons learned	Szemantikai eltérés -> közös fogalmi készlet és adatminősíté s
5. Intézményesíté s	Fenntartható képesség	Életciklus- követés; audit; folyamatos T&E; fenyegetés modell frissítése	Képességfejleszté si ciklus; folyamatos fejlesztés	Technológiai elavulás -> moduláris architektúra, folyamatos frissítés

10.3.3. Interoperabilitás növelése a NATO erőkhöz belül

A NATO interoperabilitás MI-támogatott tervezésben három szinten értelmezhető:

- (i) Technikai interoperabilitás: adatátviteli protokollok, API-k, közös biztonsági mechanizmusok (titkosítás, azonosítás, naplózás), és a több szintű biztonság (MLS) kezelése.
- (ii) Szemantikai interoperabilitás: közös adatmodellek és fogalmi készlet. A PMESII-PT, JIPOE, COA/CONOPS/OPLAN termékek fogalmai csak akkor kezelhetők géppel, ha azokat ontológiákba, taxonómiákba vagy közös adattár-sémákba rendezzük. Ennek hiányában az MI-rendszer „fordítási hibákat” fog termelni (pl. eltérően értelmezett hatásfogalom, különböző logisztikai kategóriák), amelyek a koalíciós döntéseket torzíthatják.
- (iii) Eljárásrendi interoperabilitás: a tervezési termékek előállításának és jóváhagyásának összehangolása. A koalíciós tervezésben a különböző nemzeti korlátozások (caveat), a besorolási szintek és a felelősségi körök gyakran eltérnek. Ezért az MI-rendszereknek

támogatniuk kell a „policy-aware” működést: felhasználónként és nemzetenként eltérő adat- és funkció-hozzáférést, valamint a javaslatokhoz rendelt bizonyítékokat.

A NATO digitális transzformációs céljai és az NCIA technológiai irányai javasolhatóvá teszik egy olyan architektúra kialakítását, amelyben a tervezéstámogató MI moduláris, federált módon működik: a nemzeti rendszerekben futó komponensek szövetségi szinten csak a szükséges metaadatokat és eredményeket osztják meg. Ez a megközelítés csökkenti a szenzitív adatok kitettségét, és jobban illeszkedik a koalíciós környezet realitásaihoz.[58], [76]

10.3.4. Módszertani ajánlások a mesterséges intelligencia hatékony alkalmazásához

A hatékony MI-adaptációhoz a policy és governance ajánlásoknak három célja van: (i) a kockázatok kontrollált szinten tartása, (ii) a döntéshozói felelősség és emberi kontroll biztosítása, (iii) a bevezetés gyorsítása úgy, hogy közben ne sérüljenek a jogi és etikai elvárások.

1) Risk-based bevezetés. A tervezéstámogató MI-t célszerű kockázatalapú portfólióként kezelni: külön kategóriába kerüljenek a „támogató” és a „kritikus” funkciók. Például egy dokumentum-összegző eszköz kockázata más, mint egy célkijelölésre vagy kinetikus erőalkalmazásra hatással lévő optimalizáció.

2) Traceability és audit. Kötelező követelménynek javaslom, hogy minden MI-javaslat visszavezethető legyen a bemeneti adatokra, a modellverzióra és a paraméterezésre. Ez nemcsak megfelelési, hanem műveleti tanulási követelmény is: AAR során csak így lehet megérteni, hogy egy javaslat miért működött vagy miért hibázott.

3) Felelősségi körök és emberi kontroll. A parancsnoki felelősség nem ruházható át MI-re; ezért a humán–gép csapatban a szerepeknek explicitnek kell lenniük (ki adja a bemenetet, ki hagyja jóvá a kimenetet, ki viseli a kockázatot). A human-in-the-loop nem pusztán „jóváhagyás”, hanem értelmezés és ellenőrzés.

4) T&E és red teaming. A műveleti környezetben alkalmazott MI csak folyamatos tesztelés-értékeléssel tartható megbízható szinten. Javaslom a rendszeres adversarial teszteket, a modellmérgezés és adatmanipuláció elleni próbákat, valamint a wargaming és gyakorlatok során gyűjtött teljesítményadatok beépítését.

5) Kiberbiztonság és „secure by design”. A tervezéstámogató MI-t olyan architektúrában kell bevezetni, amely támogatja a zero-trust elveket, az erős naplózást és a jogosultságkezelést. A stratégiai kiberműveletek irányelvei alapján a tervezési lánc digitális komponensei önálló támadási felületet képeznek; ezért az MI-rendszert nem lehet a C2 és információs rendszerek kiberbiztonsági kereteitől függetlenül kezelni.[80]

A fenti ajánlások módszertani gerincét a NIST AI RMF adja, amely a kormányzás, a kockázatok feltérképezése, mérése és kezelése mentén írja le a megbízható MI életciklus-szintű menedzsmentjét.[12]

10.3.5. A COPD, a PMESII adatforrások és az MI integrációja INCOSE Systems Engineering szemlélettel

A disszertáció egyik kulcsállítása, hogy a tervezéstámogató MI bevezetése a katonai szervezetekben klasszikus rendszerintegrációs feladat. Ennek megfelelően a bevezetéshez célszerű az INCOSE Systems Engineering (SE) módszertanát alkalmazni, mert az SE képes összekötni a doktrinális követelményeket, az architektúrát, a verifikációt-validációt és az életciklus-szintű kockázatkezelést.[77]

Az INCOSE SE-logika alapján a COPD-hez illesztett MI-rendszer kialakítása a következő fő lépésekre bontható:

- 1) Stakeholder igények és követelmények. A tervezők, a parancsnok, az elemzők, a jogi-etikai megfélemlőségi felelősök és a kiberbiztonsági szereplők elvárásait követelményekké kell alakítani (pl. „a COA-összevetésben a bizonytalanságok megjelenítése kötelező”).
- 2) Funkcionális bontás. A COPD termékekhez rendelt MI-funkciókat (adatgyűjtés, fúzió, elemzés, generálás, szimuláció, monitorozás) funkcionális architektúrába kell szervezni, és meg kell határozni a humán–gép csapat feladatmegosztását.
- 3) Logikai és fizikai architektúra. Meg kell tervezni a moduláris felépítést: adatcsatornák, modellek, szolgáltatások, felhasználói felületek, jogosultságkezelés, naplózás és cyber kontrollok. A koalíciós környezet miatt a federált architektúrák általában kedvezőbbek, mint a teljesen központosított megoldások.
- 4) Verifikáció és validáció (V&V). A tervezéstámogató MI-nél a V&V nemcsak technikai tesztek jelent (pontosság, megbízhatóság), hanem műveleti validációt is: a törzs felhasználói

tesztelése, gyakorlatok és wargaming során mért hasznosság, valamint megfelelőségi ellenőrzések.

5) Konfiguráció- és változásmenedzsment. A tervezési környezet változik: új adatok, új fenyegetések, új doktrinális irányok. Az MI-rendszernek ezért verziózott adat- és modellkezelésre (MLOps/ModelOps) és változáskezelésre van szüksége.

6) Kockázatkezelés és assurance. A kockázatokat (pl. bias, adversarial támadás, adatkitettség) folyamatosan mérni és kezelni kell; a compliance és audit nem utólagos dokumentáció, hanem a „digitális fonál” része.

Az SE-szemlélet fő előnye, hogy a COPD termékei és az MI-rendszer követelményei között megteremti a nyomon követhetőséget. Így a rendszer fejlesztése és üzemeltetése összhangban tartható a doktrinális változásokkal és a szövetségi digitális stratégiai irányokkal.[58], [76], [77]

10.4. Mérőszámok a mesterséges intelligencia katonai művelettervezésben történő felhasználásának értékeléséhez

A tervezéstámogató MI értékelése csak akkor hiteles, ha egyszerre vizsgálja (i) a folyamat teljesítményét (sebesség, erőforrásigény), (ii) a termékek minőségét (helyesség, konzisztencia, bizonyítékalapúság), (iii) az emberi–gép csapat működését (használhatóság, bizalom, terhelés), és (iv) a kockázatokat és megfelelőséget (biztonság, audit, robusztusság). Javasoltam ezért egy többdimenziós metrikarendszert, amely a COPD fázisaihoz illeszthető.

A mérőszámok kiválasztásánál alapelv, hogy a rendszernek „műveletileg” kell hasznosnak lennie: nem elegendő, ha egy modell laboratóriumi pontossága magas. A döntéstámogatásban az idő, a relevancia és a megbízhatóság együtt számít. A JADC2 és a szövetségi digitalizációs törekvések is arra mutatnak, hogy a jövőben a döntési ciklusidő és a közös helyzetkép minősége kulcs-képesség lesz.[58], [74]–[76]

A 10-3. táblázat a javasolt mérőszámok fő csoportjait és példáit foglalja össze.

Táblázat 10-3. Javasolt mérőszámrendszer a tervezéstámogató MI értékeléséhez

Dimenzió	Mérőszám példák	Adatforrás / mérés	COPD kapcsolódás
----------	-----------------	--------------------	------------------

Folyamat (hatékonyság)	Tervkészítési idő; ciklusidő; törzsmunkaóra; iterációk száma	Időnaplók; workflow log; gyakorlat/éles adatok	1–4. fázis: SSA, MRO/COA, CONOPS/OPLAN
Termékminőség	Konzisztencia a dokumentumok között; hivatkozások teljessége; bizonytalanság jelölése	Dokumentum-ellenőrző eszközök; reviewer pontozás	Minden fázis: tervezési termékek
Analitikai teljesítmény	Predikció kalibráltsága; FP/FN; anomália-riasztások pontossága	Ground truth (AAR); szimuláció; ellenőrzött tesztkészletek	1–2. fázis (fenyegetés); 5. fázis (monitorozás)
Robusztusság és biztonság	Adversarial teszt sikeraránya; modellmérgezés detektálás ideje; incidens-reakció idő	Red teaming; pentest; SOC naplók	Minden fázis: cyber és AML védelem
Humán-gép csapat	Felhasználói terhelés (pl. kérdőív); bizalom; elfogadottság; hibajavítási arány	User study; gyakorlat utáni értékelés; usage analytics	Törzsmunka mindvégig
Megfelelőség és governance	Audit-nyomvonal teljessége; modellverzió-követés; policy sértések száma	Governance dashboard; konfigurációmenedzsment; audit	Minden fázis: döntési pontok és jóváhagyások

A mérőszámok gyakorlati bevezetéséhez javaslom egy „tervezési MI dashboard” kialakítását, amely a fenti dimenziók mentén, de felhasználóbarát módon mutatja meg a rendszer állapotát: milyen adatforrásokra támaszkodik, milyen bizonytalansággal dolgozik, melyek a leggyakoribb hibaforrások, és milyen kockázatok aktívak. A dashboard célja, hogy a parancsnoki és törzsszintű bizalom ne narratívákon, hanem mérhető tényeken alapuljon.

Külön kiemelem, hogy a metrikák között szerepeljenek „negatív metrikák” is: például a téves riasztások költsége, a hibás javaslatok emberi javítási ideje, vagy a túlzott automatizációból

fakadó döntési torzítások. Ezek nélkül a mérőrendszer csak a sikereket méri, és nem segíti a kockázatok kontrollját.

10.5. Ajánlás további kutatási irányokra

A disszertáció több olyan kutatási irányt azonosított, amelyekben további munka szükséges ahhoz, hogy a tervezéstámogató MI képesség gyakorlati, szövetségi szinten is széles körben alkalmazható legyen. A legfontosabb javasolt irányok a következők.

1) Bizonytalanság kvantifikálása és kommunikációja. A művelettervezésben a döntések ritkán alapulnak teljes információ; ezért a predikciók és elemzések mellett a bizonytalanság megjelenítése a döntéstámogatás része. Kutatni szükséges, hogy mely UQ (uncertainty quantification) technikák és vizualizációk segítik legjobban a parancsnoki ítéletalkotást.

2) Multimodális és többnyelvű koalíciós MI. A NATO környezetben a tervezési adatok szöveges, térképi, kép- és jeladat formában, több nyelven keletkeznek. A többnyelvű terminológia- és doktrínaegységesítés MI-támogatása önálló kutatási terület, amely összekapcsolja a nyelvtechnológiát, a tudásgráfokat és a policy-aware hozzáféréskezelést.

3) Adversarial ML és ellentevékenységek elleni védelem a tervezési láncban. A kiber- és információs műveletek várhatóan egyre inkább célba veszik a döntéstámogató rendszereket. Szükséges a műveleti környezetre szabott fenyegetésmoделlek, tesztkészletek és védelmi minták kidolgozása, valamint annak vizsgálata, hogy a red teaming hogyan építhető be rutin jelleggel a gyakorlatokba.

4) „Digital twin” és kampányszintű szimulációk. A Mosaic Warfare és a distributed kill chain gondolkör azt sugallja, hogy a jövőben a moduláris képesség-összeállítás és a gyors adaptáció lesz döntő.[78], [79] Ennek megfelelően érdemes kutatni olyan digitális iker megoldásokat, amelyek kampányszinten képesek a COA-k robusztusságát, a logisztikai fenntarthatóságot és a multi-domain kölcsönhatásokat modellezni.

5) Metrikák és benchmarkok standardizálása. A 10.4. alfejezet javasolt metrikái keretrendszerként adnak, de szükséges olyan nyílt vagy ellenőrzött benchmark készletek kidolgozása, amelyekkel a tervezéstámogató MI teljesítménye összehasonlítható különböző szervezetek és eszközök között.

6) Szervezeti tanulás és tudásmenedzsment MI-vel. AAR és lessons learned folyamatok MI-támogatása (automatikus kivonatolás, ok-okozati mintázatok, ajánlások) olyan terület, amely közvetlenül javíthatja a következő tervezési ciklusok minőségét. Ennek kutatásához azonban szükséges a tanulságok strukturált rögzítése és a törzsszintű tudásgráfok kialakítása.

Összességében a jövőbeli kutatásnak azt a hidat kell tovább erősítenie, amely a doktrinális tervezési folyamatok (COPD), a modern adat- és digitális stratégiák (JADC2, NATO digital transformation), valamint a megbízható MI mérnöki és governance keretei között húzódik.[58], [74]–[76], [12], [77]

IRODALOMJEGYZÉK

- [1] NATO, “NATO 2022 Strategic Concept”, 29 June 2022, pp. 3–5 (PDF p. 3–5). Elérhető: <https://www.nato.int/strategic-concept/>
- [2] European External Action Service (EEAS), “A Strategic Compass for Security and Defence”, March 2022, pp. 4–7. Elérhető: https://www.eeas.europa.eu/eeas/strategic-compass-security-and-defence-1_en
- [3] Magyarország Kormánya, „1163/2020. (IV. 21.) Korm. határozat Magyarország Nemzeti Biztonsági Stratégiájáról – ‘Biztonságos Magyarország egy változékony világban’”, Magyar Közlöny 2020/81, pp. 1–2 (technológiai forradalom; technológiai/kibertérbeli biztonság). Elérhető: https://nbsz.gov.hu/docs/NBS_MK_2020_81_1163.2020_Korm.hat.pdf
- [4] NATO Allied Command Operations, “Comprehensive Operations Planning Directive (COPD), Version 3.1”, NATO UNCLASSIFIED, 2023, pp. 29–30 (PDF p. 29–30).
- [5] „A NATO átfogó művelettervezési útmutatójának evolúciója (COPD)”, “Hadtudomány” (tanulmány), 2023.02.07., pp. 2–3 és 9–10 (PDF oldalszámok).
- [6] NATO, “NATO’s Digital Transformation Implementation Strategy”, 17 Oct. 2024, releváns szakaszok (interoperabilitás, MDO, felelős MI/adat). Elérhető: https://www.nato.int/cps/en/natohq/topics_238776.htm
- [7] NATO, “Data Strategy for the Alliance”, 05 May 2025, releváns szakaszok (adatmegosztás, adat-kormányzás). Elérhető: https://www.nato.int/cps/en/natohq/topics_272784.htm
- [8] NATO, “Allied Joint Doctrine for the Planning of Operations (AJP-5), Edition A Version 2 + UK national elements (Change 1)”, 10 Mar. 2021, p. 3-3 (operations design mint iteratív megértés; PDF p. 55). Elérhető: https://assets.publishing.service.gov.uk/media/6054d017e90e0724be025a8f/20210310-AJP_5_with_UK_elem_final_web.pdf
- [9] U.S. Joint Chiefs of Staff, “Joint Planning (JP 5-0)”, 01 Dec. 2020, pp. III-5–III-6 (iteratív párbeszéd a parancsnok és törzs között; PDF p. 83–84). Elérhető: <https://www.safety.marines.mil/Portals/92/Ground%20Safety%20for%20Marines%20%28GSM%29/References%20Tab/JP%205-0%20Joint%20Planning%20PDF.pdf>
- [10] U.S. Joint Chiefs of Staff, “Joint Intelligence Preparation of the Operational Environment (JP 2-01.3)”, 21 May 2014, p. I-1 (JIPOE definíció és 4 lépés; PDF p. 20). Elérhető: https://usna.edu/Training/_files/documents/References/2C%20MQS%20References/2015-2016%20MQS/Joint%20Publication%20JP%202-01.3%20Joint%20Intelligence%20Preparation%20of%20the%20Operating%20Environment.pdf
- [11] Magyarország Kormánya, „1393/2021 (VI. 24.) Korm. határozat Magyarország Nemzeti Katonai Stratégiájáról” (Honvédelem portál), 2021.12.04., bevezető rész. Elérhető: <https://honvedelem.hu/hirek/nemzeti-katonai-strategia.html>
- [12] National Institute of Standards and Technology (NIST), “Artificial Intelligence Risk Management Framework (AI RMF 1.0)”, NIST AI 100-1, Jan. 2023, pp. 1–3 és 20–22 (GOVERN–MAP–MEASURE–MANAGE). Elérhető: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>

- [13] U.S. Department of Defense, “Responsible Artificial Intelligence Strategy and Implementation Pathway,” Jun 2024 (RAI S&I Pathway). Elérhető: <https://media.defense.gov/2024/Oct/26/2003571790/-1/-1/0/2024-06-RAI-STRATEGY-IMPLEMENTATION-PATHWAY.PDF> (letöltve: 2025-12-24). (Hivatkozott oldalak: pp. 1–2).
- [14] NATO, “Emerging and disruptive technologies,” NATO Topic (updated 25 Jun. 2025). Elérhető: <https://www.nato.int/en/what-we-do/deterrence-and-defence/emerging-and-disruptive-technologies> (letöltve: 2025-12-24).
- [15] NATO, “NATO’s Data and Artificial Intelligence Review Board,” NATO Official text, 13 Oct. 2022. Elérhető: <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2022/10/13/natos-data-and-artificial-intelligence-review-board> (letöltve: 2025-12-24).
- [16] NATO, “Summary of NATO’s revised Artificial Intelligence (AI) strategy,” NATO Official text, 10 Jul. 2024. Elérhető: https://www.nato.int/cps/en/natohq/official_texts_227237.htm (letöltve: 2025-12-24).
- [17] Négyesi–Szabadszöke, “Military Strategic OPLAN supported by AI” (kézirat / szerzői PDF), PDF hivatkozott oldal: p. 30. (A disszertációhoz a szerzők által rendelkezésre bocsátott változat).
- [18] NATO Allied Command Transformation, “Fact Sheet: Military Uses of Artificial Intelligence, Automation, and Robotics (MUAAR),” May 2019, p. 1. Elérhető: https://www.act.nato.int/wp-content/uploads/2023/05/2019_MCDC_MUAAR-2.pdf (letöltve: 2025-12-24).
- [19] U.S. Department of Defense, DoD Directive 3000.09, “Autonomy in Weapon Systems,” 25 Jan 2023. Elérhető: <https://www.esd.whs.mil/portals/54/documents/dd/issuances/dodd/300009p.pdf> (letöltve: 2025-12-24). (Hivatkozott oldalak: Glossary/Definitions, pp. 21–23; továbbá a „failure” okai, pp. 22–23).
- [20] NATO, “Summary of NATO’s Autonomy Implementation Plan,” NATO Official text, 13 Oct. 2022. Elérhető: <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2022/10/13/summary-of-natos-autonomy-implementation-plan> (letöltve: 2025-12-24).
- [21] U.S. Department of Defense, “Implementing Responsible Artificial Intelligence in the Department of Defense,” Memorandum, 26 May 2021. Elérhető: <https://media.defense.gov/2021/May/27/2002730593/-1/-1/0/IMPLEMENTING-RESPONSIBLE-ARTIFICIAL-INTELLIGENCE-IN-THE-DEPARTMENT-OF-DEFENSE.PDF> (letöltve: 2025-12-24). (Hivatkozott oldalak: pp. 1–2).
- [22] U.S. Department of Defense, “Ethical Principles for Artificial Intelligence,” Appendix B, 2020. Elérhető: <https://imlive.s3.amazonaws.com/Federal%20Government/ID151922140063916450446479915674517979637/Attachment%200005%20Appendix%20B%20Department%20of%20Defense%20AI%20Ethical%20Principles.pdf> (letöltve: 2025-12-24).
- [23] NATO, “Summary of the NATO Artificial Intelligence Strategy,” NATO Official text, 22 Oct. 2021. Elérhető: <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2021/10/22/summary-of-the-nato-artificial-intelligence-strategy> (letöltve: 2025-12-24).
- [24] NIST, “Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and Mitigations (NIST AI 100-2e2023),” 2023. Elérhető: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-2e2023.pdf> (letöltve: 2025-12-24). (Hivatkozott oldalak: pp. 9–16, különösen a data poisoning/evasion fogalmak).
- [25] MITRE, “ATLAS (Adversarial Threat Landscape for Artificial-Intelligence Systems),” weboldal. Elérhető: <https://atlas.mitre.org/> (letöltve: 2025-12-24).

- [26] OpenSSF, "Visualizing Secure MLOps (MLSecOps): A Practical Guide for Building Robust AI/ML Pipeline Security," Whitepaper, Aug 2025. Elérhető: https://openssf.org/wp-content/uploads/2025/08/OpenSSF_MLSecOps_Whitepaper.pdf (letöltve: 2025-12-24). (Hivatkozott oldalak: pp. 3, 26–28, 40).
- [27] NATO Standardization Office, "Allied Joint Doctrine for the Planning of Operations (AJP-5)," Edition A Version 2, 2019. [Online]. Elérhető: https://www.coemed.org/files/stanags/01_AJP/AJP-5_EDA_V2_E_2526.pdf (letöltve: 2025-12-24), pp. 16–17, 23, 129–130.
- [28] Sun Tzu, "The Art of War" (ford. L. Giles), 1910. [Online]. Elérhető: <https://ia600502.us.archive.org/12/items/TheArtOfWarBySunTzu/ArtOfWar.pdf> (letöltve: 2025-12-24), pp. 31–32, 42.
- [29] M. D. Krause és R. C. Phillips (szerk.), "Historical Perspectives of the Operational Art," U.S. Army Center of Military History, 2005. [Online]. Elérhető: <https://history.army.mil/portals/143/Images/Publications/catalog/70-89-1.pdf> (letöltve: 2025-12-24), pp. 4–6.
- [30] NATO Comprehensive Operational Planning Directive (COPD), v3.1, 2023. (PDF)
- [31] F. Fazekas, "A NATO átfogó művelettervezési irányelve (COPD) evolúciója" (tanulmány, PDF). pp. 1, 9.
- [32] Chairman of the Joint Chiefs of Staff, "Joint Publication 5-0: Joint Planning," 1 Dec 2020. (PDF p. 33.)
- [33] NATO Allied Command Operations, Comprehensive Operations Planning Directive (COPD) v3.1, NATO, Oct. 2023.
- [34] NATO, Allied Joint Publication AJP-5: Allied Joint Doctrine for the Planning of Operations, Ed. A, Ver. 2, with UK Change 1, Mar. 2021. [Online]. Available: https://assets.publishing.service.gov.uk/government/uploads/system/uploads/attachment_data/file/974694/doctrine_uk_ajp_5.pdf (Accessed: 2025-12-23).
- [35] O. Jurevicius, "PMESII-PT explained," Strategic Management Insight, Jun. 16, 2025. [Online]. Available: <https://strategicmanagementinsight.com/tools/pmesii-pt.html> (accessed: 2025-12-23).
- [36] UK Ministry of Defence, Intelligence, Counter-Intelligence and Security Support to Joint Operations (JDP 2-00), 4th ed., Aug. 2023. [Online]. Available: https://assets.publishing.service.gov.uk/media/653a4b0780884d0013f71bb0/JDP_2_00_Ed_4_web.pdf (accessed: 2025-12-23).
- [37] Office of the Director of National Intelligence, Intelligence Community Directive (ICD) 203: Analytic Standards, Jan. 2015. [Online]. Available: <https://www.dni.gov/files/documents/ICD/ICD%20203%20Analytic%20Standards.pdf> (accessed: 2025-12-23).
- [38] U.S. Department of Defense, Joint and National Intelligence Support to Military Operations (JP 2-01), Jul. 5, 2017. [Online]. Available: https://usna.edu/Training/_files/jp2_01_20170705v2.pdf (accessed: 2025-12-23).
- [39] C. von Sikorski and M. Hameleers, "Dezinformáció a mesterséges intelligencia (MI) korában: következmények az újságírásra és a tömegkommunikációra nézve," *Journalism & Mass Communication Quarterly*, vol. 102, no. 4, first published online Nov. 30, 2025, doi: 10.1177/10776990251375097.

- [40] B. B. Jensen, *The Future of Intelligence Analysis: Technology, People and Processes*. London, UK: Routledge, 2021.
- [41] J. R. Clapper, *Facts and Fears: Hard Truths from a Life in Intelligence*. New York, NY, USA: Viking, 2018.
- [42] U.S. Department of Defense, Joint Planning (JP 5-0), Dec. 1, 2020. [Online]. Available: https://irp.fas.org/doddir/dod/jp5_0.pdf (accessed: 2025-12-23).
- [43] K. M. Sayler, *Artificial Intelligence and National Security*, Congressional Research Service Report R45178, 2020.
- [44] National Security Commission on Artificial Intelligence, *Final Report*, Mar. 2021. [Online]. Available: <https://www.nscai.gov/wp-content/uploads/2021/03/Full-Report-Digital-1.pdf> (accessed: 2025-12-23).
- [45] NATO, "NATO releases first-ever strategy for Artificial Intelligence," Oct. 2021. [Online]. Available: https://www.nato.int/cps/en/natohq/news_187902.htm (accessed: 2025-12-23).
- [46] J. Schubert, J. Brynielsson, M. Nilsson, and P. Svenmarck, "Artificial Intelligence for Decision Support in Command and Control Systems," in *Proc. Int. Conf. on Information Fusion*, 2018.
- [47] P. Svenmarck, L. Luotsinen, M. Nilsson, and J. Schubert, "Possibilities and Challenges for Artificial Intelligence in Military Applications," in *Proc. Int. Conf. on Information Fusion*, 2018.
- [48] R. Ehlers and P. Blannin, "Integrated Planning and Campaigning for Complex Problems," *Parameters*, vol. 51, no. 2, 2021. Available: <https://press.armywarcollege.edu/parameters/vol51/iss2/10> (accessed: 2025. 12. 28.).
- [49] D. Wilkins and R. Desimone, "Applying an AI Planner to Military Operations Planning," 1996.
- [50] A. Tate, J. Levine, P. Jarvis and J. Dalton, "Using AI Planning Technology for Army Small Unit Operations," 2000.
- [51] SRI International, "SIPE-2: System for Interactive Planning and Execution." Available: <http://www.ai.sri.com/~sipe/> (accessed: 2025. 12. 28.).
- [52] T. M. Kinyon, "Validating a Social Model Wargame: An Analysis of the Green Country Model," Naval Postgraduate School, 2012. Available: <https://www.hsdl.org/?abstract&did=732136> (accessed: 2025. 12. 28.).
- [53] A. James, J. Pierowicz, M. D. Moskal, T. Hanratty, D. Tuttle, B. Sensenig and B. Hedges, "Assessing Consequential Scenarios in a Complex Operational Environment Using Agent-Based Simulation," ARL-TR-7954, U.S. Army Research Laboratory, 2017. Available: <https://www.hsdl.org/?abstract&did=817446> (accessed: 2025. 12. 28.).
- [54] Magyar Honvédség, "A Magyar Honvédség Törzsszolgálati szakutasítása," Ált-4/457, 2013. (PDF, felt
- [55] U.S. Army War College, "Strategic Cyberspace Operations Primer," 18 Dec. 2023. (PDF, 2025. 12. 28.).
- [56] NATO Allied Command Operations (ACO), *Comprehensive Operations Planning Directive (COPD)*, Ver. 3.1, belső kiadvány, n.d.
- [57] NATO, *NATO 2022 Strategic Concept*, Jun. 29, 2022. [Online]. Available: <https://www.act.nato.int/wp-content/uploads/2023/05/290622-strategic-concept.pdf> (Accessed: 2025-12-23).

- [58] NATO, "NATO's Digital Transformation Implementation Strategy," Oct. 17, 2024. [Online]. Available: <https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2024/10/17/natos-digital-transformation-implementation-strategy> (Accessed: 2025-12-23).
- [59] U.S. Department of Defense, JP 2-0: Joint Intelligence, Oct. 22, 2013. [Online]. Available: https://irp.fas.org/doddir/dod/jp2_0.pdf (Accessed: 2025-12-23).
- [60] Office of the Director of National Intelligence, Intelligence Community Directive (ICD) 203: Analytic Standards, Jan. 2015. [Online]. Available: <https://www.dni.gov/files/documents/ICD/ICD-203.pdf> (Accessed: 2025-12-23).
- [61] B. B. Jensen, *The Future of Intelligence Analysis: Technology, People and Processes*. Abingdon, UK: Routledge, 2021. [Online]. Available: <https://www.routledge.com/The-Future-of-Intelligence-Analysis/Jensen/p/book/9780367722326> (Accessed: 2025-12-23).
- [62] NATO Science & Technology Organization, *Science & Technology Trends 2020–2040: Exploring the S&T Edge*, 2020. [Online]. Available: <https://www.sto.nato.int/document/science-technology-trends-2020-2040-exploring-the-st-edge/> (Accessed: 2025-12-23).
- [63] J. R. Clapper, *Facts and Fears: Hard Truths from a Life in Intelligence*. New York, NY, USA: Viking, 2018. [Online]. Available: <https://www.penguinrandomhouse.com/books/606574/facts-and-fears-by-james-r-clapper/> (Accessed: 2025-12-23).
- [64] European Parliament and Council of the European Union, Regulation (EU) 2024/1689 of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union acts, 2024. [Online]. Available: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng> (Accessed: 2025-12-23).
- [65] National Security Commission on Artificial Intelligence, *Final Report*, Mar. 2021. [Online]. Available: <https://www.dwt.com/-/media/files/blogs/artificial-intelligence-law-advisor/2021/03/nscai-final-report--2021.pdf> (Accessed: 2025-12-23).
- [66] National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1)*, Jul. 26, 2024. [Online]. Available: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf> (Accessed: 2025-12-23).
- [67] NATO, "Summary of the NATO Artificial Intelligence Strategy," Oct. 22, 2021. [Online]. Available: https://www.nato.int/cps/en/natohq/news_185064.htm?selectedLocale=en (Accessed: 2025-12-23).
- [68] P. Lin, K. Abney, and R. Jenkins, Eds., *Robot Ethics 2.0: From Autonomous Cars to Artificial Intelligence*. Oxford University Press, 2017. [Online]. Available: <https://global.oup.com/academic/product/robot-ethics-20-9780190652951> (Accessed: 2025-12-23).
- [69] International Committee of the Red Cross, *Autonomous Weapon Systems: Implications of Increasing Autonomy in the Critical Functions of Weapons*, Sep. 1, 2017. [Online]. Available: <https://www.icrc.org/en/publication/4283-autonomous-weapons-systems> (Accessed: 2025-12-23).
- [70] United Nations Office for Disarmament Affairs, "Group of Governmental Experts on Lethal Autonomous Weapons Systems (GGE LAWS)," 2024. [Online]. Available: <https://meetings.unoda.org/meeting/ccwgge-laws-2024/> (Accessed: 2025-12-23).
- [71] A. N. Brundage et al., "Toward Trustworthy AI Development: Mechanisms for Adversarial ML," arXiv:2006.07559, 2020. [Online]. Available: <https://arxiv.org/abs/2006.07559> (Accessed: 2025-12-23).

- [72] I. Goodfellow, J. Shlens, and C. Szegedy, "Explaining and Harnessing Adversarial Examples," arXiv:1412.6572, 2014. [Online]. Available: <https://arxiv.org/abs/1412.6572> (Accessed: 2025-12-23).
- [73] NATO Allied Command Operations (ACO), "Comprehensive Operational Planning Directive (COPD) v3.1," internal copy, 2021.
- [74] U.S. Department of Defense, "Summary of the Joint All-Domain Command and Control (JADC2) Strategy," 2022. [Online]. Available: <https://media.defense.gov/2022/Mar/17/2002958406/-1/-1/1/SUMMARY-OF-THE-JOINT-ALL-DOMAIN-COMMAND-AND-CONTROL-STRATEGY.PDF> (accessed: 2025-12-23).
- [75] U.S. Government Accountability Office, "Battle Management: DOD Is in Early Stages of Defining JADC2," GAO-23-105495, 2023. [Online]. Available: <https://www.gao.gov/assets/gao-23-105495.pdf> (accessed: 2025-12-23).
- [76] NATO Communications and Information Agency (NCIA), "NCIA Technology Strategy," 2023/2024. [Online]. Available: https://www.ncia.nato.int/resources/site1/General/newsroom/publications/Public_NCIA_Technology%20Strategy_external_v6%20-%20digital.pdf (accessed: 2025-12-23).
- [77] International Council on Systems Engineering (INCOSSE), Systems Engineering Handbook, 5th ed., 2023. [Online]. Available: <https://www.incose.org/products-and-publications/se-handbook> (accessed: 2025-12-23).
- [78] Defense Advanced Research Projects Agency (DARPA), "Mosaic Warfare," [Online]. Available: <https://www.darpa.mil/news/mosaic-warfare> (accessed: 2025-12-23).
- [79] RAND Corporation, "Distributed Kill Chains: Drawing Insights for Mosaic Warfare from the Immune System and Other Complex Adaptive Systems," 2021. [Online]. Available: https://www.rand.org/content/dam/rand/pubs/research_reports/RR500/RR573-1/RAND_RRA573-1.pdf (accessed: 2025-12-23).
- [80] U.S. Army War College, "Strategic Cyberspace Operations Guide," 2023 update. [Online]. Available: https://media.defense.gov/2023/Oct/02/2003312499/-1/-1/0/STRATEGIC_CYBERSPACE_OPERATIONS_GUIDE.PDF (accessed: 2025-12-23).
- [81] NATO Science & Technology Organization (STO), "Science & Technology Trends: 2020–2040 – Exploring the S&T Edge," 2020. Elérhető: https://securitydelta.nl/media/com_hsd/report/406/document/190422-ST-Tech-Trends-Report-2020-2040.pdf (hivatkozott oldalak: pp. 1–6, 9–13, különösen p. 9 és p. 12).
- [82] European Commission, "Annexes 1 to 9 to the Proposal for a Regulation ... Laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act)," COM(2021) 206 final, Annex I, p. 2, 21 Apr. 2021. Elérhető: <https://artificialintelligenceact.eu/wp-content/uploads/2022/05/AIA-COM-Annexes-21-April-21.pdf> (letöltve: 2025-12-24).
- [83] Szabadföldi István, "Artificial Intelligence in Military Application – Opportunities and Challenges," Land Forces Academy Review, vol. 26, no. 2, 2021, pp. 157–165 (különösen pp. 161–164). Elérhető: https://www.armyacademy.ro/reviste/2_2021/ART_11.pdf (letöltve: 2025-12-24).
- [84] Szabadföldi István, "A mesterséges intelligenciával támogatott nyílt információszerezés (OSINT) – evolúció és kihívások," Nemzetbiztonsági Szemle, 10. évfolyam, 1. szám, 2022, pp. 30–51. Elérhető: <https://real.mtak.hu/156537/> (letöltve: 2025-12-24).
- [85] High-Level Expert Group on Artificial Intelligence (European Commission), "Ethics Guidelines for Trustworthy AI," 2019. Elérhető:

https://www.europarl.europa.eu/meetdocs/2014_2019/documents/libe/dv/ai_guidelines/ai_guidelines_en.pdf (letöltve: 2025-12-24). (Hivatkozott oldalak: pp. 13–20, különösen p. 14).

[86] „• AI application in Electronic Warfare, 4th Annual International Research Conference GlobState 2021, Poland, <https://cdissz.wp.mil.pl/en/articlesnews-2022/globstate-2021-monograph/>

[87] C. von Clausewitz, "On War" (ford. J. J. Graham), 1873. [Online]. Elérhető: https://www.files.ethz.ch/isn/125489/5003_OnWar.pdf (letöltve: 2025-12-24), pp. 11–12.

[88] Allied Command Operations (ACO), Comprehensive Operational Planning Directive (COPD), Version 3.1, NATO UNCLASSIFIED, 13 Oct. 2023. (Belső dokumentum; példa oldalak: fázislogika pp. 76–83; accelerated planning pp. 98–101; decision briefs/handover pp. 103–108; kockázat pp. 92–96; assessment pp. 172–186.)

[89] UK Ministry of Defence, Allied Joint Doctrine for the Planning of Operations (AJP-5), Edition A, Version 2, 2021. Elérhető: <https://www.gov.uk/government/publications/allied-joint-doctrine-for-the-planning-of-operations-ajp-5> (letöltve: 2025-12-24).

[90] NATO Joint Analysis and Lessons Learned Centre (JALLC), The NATO Lessons Learned Handbook, 4th ed., Jun. 2022. Elérhető: https://nllp.jallc.nato.int/iks/sharing%20public/jallc_ll_handbook_update_-_4th_edition_final_14072022.pdf (letöltve: 2025-12-24).

[91] Joint Staff J-7, Commander's Handbook for Assessment Planning and Execution, Version 1.0, Sep. 2011. Elérhető: https://www.jcs.mil/Portals/36/Documents/Doctrine/pams_hands/assessment_hbk.pdf (letöltve: 2025-12-24).

[92] Joint Chiefs of Staff, Joint Publication 5-0: Joint Planning, 16 Jun. 2017. Elérhető: https://fas.org/irp/doddir/dod/jp5_0.pdf (letöltve: 2025-12-24).

[93] J. Strange, Centers of Gravity & Critical Vulnerabilities: Building on the Clausewitzian Foundation so that We Can All Speak the Same Language, Marine Corps University, 1996. Elérhető: [https://commons.wikimedia.org/wiki/File:Centers_of_gravity_%26_critical_vulnerabilities_-_building_on_the_Clausewitzian_foundation_so_that_we_can_all_speak_the_same_language_\(IA_centersofgravit00stra\).pdf](https://commons.wikimedia.org/wiki/File:Centers_of_gravity_%26_critical_vulnerabilities_-_building_on_the_Clausewitzian_foundation_so_that_we_can_all_speak_the_same_language_(IA_centersofgravit00stra).pdf) (letöltve: 2025-12-24).

[94] NATO Terminology Office, NATO Terminology Database (NATOTerm), online database. Elérhető: <https://nso.nato.int/natoterm/> (letöltve: 2025-12-24).

[95] NATO, "Strategic Concept 2022," Madrid Summit, 29 Jun 2022. Elérhető: <https://www.nato.int/content/dam/nato/webready/documents/publications-and-reports/strategic-concepts/2022/290622-strategic-concept.pdf> (letöltve: 2025-12-24), pp. 1–5.

[96] NATO Standardization Office, "NATO Standard AJP-01: Allied Joint Doctrine," (AJP-01, STANAG 2437). Elérhető: https://www.coemed.org/files/stanags/01_AJP/AJP-01_EDE_V1_E_2437.pdf (letöltve: 2025-12-24), különösen Chapter 2, pp. 2–6–2–10 (PDF pp. ~66–70).

[97] Council of the European Union, "Council Conclusions on the Integrated Approach to External Conflicts and Crises," ST 5413/2018 INIT, 22 Jan 2018. Elérhető: <https://data.consilium.europa.eu/doc/document/ST-5413-2018-INIT/en/pdf> (letöltve: 2025-12-24), pp. 1–4.

[98] NATO, “NATO Operations Assessment Handbook,” 01 Jul 2015. Elérhető: https://archive.org/download/NATOOperationsAssessmentHandbook2015/NATO_Operations_Assessment_Handbook_2015.pdf (letöltve: 2025-12-24), MOP/MOE és assessment folyamat: PDF pp. 15–18, 35–40.

[99] Council of the European Union, “EU Concept for Military Planning at the Political and Strategic Level,” ST 6432/2015 INIT, 19 Mar 2015. Elérhető: <https://data.consilium.europa.eu/doc/document/ST-6432-2015-INIT/en/pdf> (letöltve: 2025-12-24), CONOPS/OPLAN logika: PDF pp. 11–12.

[100] Council of the European Union, “EU Concept for Military Command and Control,” ST 8798/2019 INIT, 13 May 2019. Elérhető: <https://data.consilium.europa.eu/doc/document/ST-8798-2019-INIT/en/pdf> (letöltve: 2025-12-24), MPCC és OpCdr szerepek: PDF pp. 7–12, 17–20.

[101] Council of the European Union, “Council Decision (CFSP) 2017/971 of 8 June 2017,” Official Journal L 146/133. Elérhető: <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX%3A32017D0971> (letöltve: 2025-12-24), különösen Article 1–3, PDF pp. 1–3.

[102] Magyar Honvédség, “Összhaderőnemi Doktrína (Ált/44),” (PDF). NATO/AJP kapcsolódások: PDF pp. 7–12, 62.

PUBLIKÁCIÓS JEGYZÉK

<https://m2.mtmt.hu/gui2/?type=authors&mode=browse&sel=10074278>

- **Financial Centers-a regional chance for Budapest?** Műhelytanulmányok, 1995, Állami Biztosításfelügyelet <https://www.jstor.org/stable/i40071540>
- <https://www.jstor.org/stable/41468264>
- **Development of the Hungarian capital market**
Central European Management Review, Prague 1996
- **Budapest as Financial Center**
Vezetéstudomány, 1996
- **Az informatikai fejlesztések problémái a bankszektorban**
Bankvilág, 1997 Október
- **Budapest mint KKE regionális üzleti központja-esélyek és erőfeszítések**
Bankvilág, 1998 Június
- **Regional Corporate Survey**
Gazdasági Minisztérium - Tanulmány Bancraft Kft. 2001
- **Final Report for the Conference on Budapest as a Regional Business Center**
Gazdasági Minisztérium Kiadvány 2002
https://drive.google.com/drive/folders/1kBJZy2ejYcHa3YB9_VnQKR4-Joq2eoU?usp=sharing
- **Budapest International Business Center Corporate Survey -**
Nemzetgazdasági Minisztérium – Bancraft Tanulmány 2013
https://drive.google.com/drive/folders/1kBJZy2ejYcHa3YB9_VnQKR4-Joq2eoU?usp=sharing
- **Artificial Intelligence in Military Application – Opportunities and Challenges**
REVISTA ACADEMIEI FORTELOR TERESTRE / LAND FORCES ACADEMY REVIEW (2247-840X 1582-6384): 26 2 pp 157-165 (2021)
DOI: <https://doi.org/10.2478/raft-2021-0022>
- **AI application in Electronic Warfare**
4th Annual International Research Conference GlobState 2021, Poland
<https://cdissz.wp.mil.pl/en/articlesnews-2022/globstate-2021-monograph/>

- **A mesterséges intelligenciával támogatott nyílt információszerzés (OSINT) – evolúció és kihívások** (NKE-KMDI The Artificial Intelligence supported OSINT – evolution and challenges)
Nemzetbiztonsági Szemle - Nationa Security Review, Budapest, NUPS, 2021
<https://folyoirat.ludovika.hu/index.php/nbsz/issue/view/460>
- **A mesterséges intelligencia alkalmazási lehetőségei az elektronikai hadviselésben**
Ludovika kiadó; „Szemelvények a katonai Műszaki tudományok eredményeiből III”
<https://webshop.ludovika.hu/termek/konyvek/hadtudomany/szemelvenyek-a-katonai-muszaki-tudomanyok-eredmenyeibol-iii/>
- **The Comprehensive Approach of Military Operation Planning and its support by Artificial Intelligence**; Col. Imre Négyesi & István Szabadföldi, Hungarian Defence Review, 2023;
<https://kiadvany.magyarhonvedseg.hu/index.php/honvszemle/issue/view/143/141>
- **Ukrajna – nemzetállam vagy orosz birodalmi kerület** (Ukraine - nation state or Russian imperial district) 2022; NKE Biztonságpolitika.hu szakportál;
<https://biztonsagpolitika.hu/egyeb/ukrajna-nemzetallam-vagy-orosz-birodalmi-kerulet?fbclid=IwAR1naOFZvVI17d1zyVN8ikAA2ACPTxuRET67Ek0HdyvPIWIIOW-KrCAmCc8>
- **Platform as a Service application in Military Cloud Architecture** (from monolithic applications to a microservices architecture); American Journal of Research, Education and Development
ISSN 2471-9986; <https://red.devlart.hu/>
- **Status Quo turning point**, University of Public Service, Hungary, Security Policy Portal 2025, <https://biztonsagpolitika.hu/in-english/analyses-in-english/status-quo-turning-point>

MELLÉKLETEK

Melléklet A – Tervezési sablonok és kitölthető űrlapok (COPD-kompatibilis)

A melléklet célja, hogy a fejezetben tárgyalt módszertani elemeket gyakorlati, kitölthető sablonok formájában összefoglalja. A sablonok nem helyettesítik a NATO hivatalos formátumait, de a COPD logikájával kompatibilis mezőstruktúrát adnak. A sablonok különösen gyorsított tervezésben hasznosak, mert biztosítják a „döntési minimum” dokumentálását: kérdés, opciók, feltételezések, kockázat, indikátorok és döntési triggerek.

A.1. Döntési napló (Decision Log) – sablon

Kitöltési útmutató: minden parancsnoki döntést (DP) egy sorban kell rögzíteni, hivatkozva a kapcsolódó termékre (SSA/MRO/CONOPS/OPLAN/FRAGO).

Dátum/Idő	DP	Döntési kérdés	Opciók	Döntés	Indok	Kockázat/mitigáció

Melléklet A.1. Döntési napló sablon (saját szerkesztés).

A.2. Feltételezéslista (Assumption Register) – sablon

Kitöltési útmutató: a feltételezés akkor elfogadható, ha van felelős és tervezett validációs pont (mikor és miből ellenőrizzük).

ID	Feltételezés	Miért szükséges?	Validáció (forrás/idő)	Owner	Státusz (nyitott/igazolt/megkérdőjelezett)

Melléklet A.2. Feltételezéslista sablon (saját szerkesztés).

A.3. Szinkronizációs mátrix (Synchronization Matrix) – egyszerűsített sablon

Kitöltési útmutató: LoE/LoO mentén, fázisonként rögzítendőek a kulcsfeladatok és a döntési pontok.

LoE/LoO	Fázis 0/1	Fázis 2	Fázis 3	Fázis 4	Fázis 5+
---------	-----------	---------	---------	---------	----------

Melléklet A.3. Szinkronizációs mátrix sablon (saját szerkesztés).

A.4. Indikátor-katalógus (MOP/MOE) – sablon

Kitöltési útmutató: minden DC-hez legfeljebb néhány kiemelt MOE-t és vezető indikátort érdemes rendelni, küszöbökkel.

DC/OO	Indikátor	Típus (MOP/MOE)	Leading/Lagging	Adatforrás	Gyakoriság	Küszöb	Felelős

Melléklet A.4. Indikátor-katalógus sablon (saját szerkesztés).

A.5. Handover/átadás ellenőrzőlista (Shift handover checklist)

Kitöltési útmutató: váltáskor röviden rögzíteni kell a döntési pontok állapotát, CCIR-t és kockázatokat.

Tétel	Megjegyzés
Jelenlegi helyzetkép változása (last 24h)	
Aktív döntési pontok és közelgő triggererek	
Nyitott CCIR és várható beérkezési idő	
Kiemelt kockázatok (Top 5) és mitigáció státusz	
FRAGO-k / változások a tervben (baseline vs current)	
Koalíciós/partner korlátozások, caveat változások	
Következő parancsnoki brief időpontja és agenda	

Melléklet A.5. Handover ellenőrzőlista (saját szerkesztés).

Melléklet B – SSA (Strategic Assessment) részletes szerkezete és kitöltési útmutató

A COPD-ben az SSA nem egyszerű „helyzetjelentés”, hanem a problémakeretezés döntéstámogató dokumentuma. Az SSA minősége meghatározza a további fázisok minőségét: ha az end state, a success criteria, a korlátozások és a feltételezések nincsenek egyértelműen rögzítve, a COA-fejlesztés során a törzs eltérő problémaképek alapján dolgozik, és a wargaming eredményei nehezen lesznek összevethetők. A melléklet egy COPD-kompatibilis SSA-vázlatot ad, fejezetenként rövid kitöltési szabályokkal és gyakori hibákkal.

B.1. Executive summary és döntési kérdés

Az SSA első oldalának feladata a döntési kérdés rögzítése. A jó executive summary három elemet tartalmaz: (1) mi történt és miért releváns; (2) mi a műveleti probléma (problem statement) és a kívánt stratégiai végállapot; (3) milyen azonnali döntés vagy iránymutatás szükséges a parancsnoktól. Fontos, hogy a summary ne legyen „információlista”: a parancsnok itt nem részleteket keres, hanem a problémakeret logikáját. Az SSA-ban a döntési kérdés változhat a helyzet pontosodásával, de minden változást verzióban és döntési naplóban rögzíteni kell, különben a traceability sérül.

- Tipikus hiba: túl hosszú summary, amely elveszíti a fókuszot és nem vezet döntéshez.
- Tipikus hiba: a problem statement csak tüneteket sorol (például incidensek), nem pedig ok–okozati viszonyt.
- Jó gyakorlat: a summary végén 3–5 legfontosabb „ismeretlen” (gap) kiemelése és owner kijelölése.

B.2. Mandátum, felhatalmazás, politikai célok és korlátok

A műveleti tervezés egyik leggyakoribb kudarca, hogy a törzs a katonai megoldás optimumát keresi anélkül, hogy a politikai és jogi kereteket a tervezés elején rögzítené. A COPD-ben a mandátum és a politikai célok rögzítése azért kritikus, mert a későbbi COA-k elfogadhatóságát (acceptability) alapvetően ez határozza meg. A korlátok (constraints) és korlátozások (restraints) különbsége praktikus: constraints azok, amiket meg kell tenni (például védelem biztosítása), restraints azok, amiket nem szabad (például egyes célok tiltása). Többnemzeti műveletekben ide tartozik a nemzeti caveat-ek összegzése is.

A kitöltéskor célszerű a korlátokat „hatásmechanizmus” szerint csoportosítani: hogyan befolyásolják az erőalkalmazást, a kommunikációt, a logisztikai útvonalakat vagy a célkijelölést. Így a korlát nem a dokumentum végén „megúszós” lista lesz, hanem a COA-fejlesztés inputja. Ha a korlátok változnak (például ROE szigorodik), akkor az SSA frissítése

mellett a risk register és a DST/DSM is módosulhat, mert a döntési pontok és triggerok más logikát kapnak.

B.3. Műveleti környezet (OE) – rendszerleírás

A COPD az OE-t rendszerként kezeli, nem pedig különálló „dimenziók” összességéként. Gyakorlatban ez azt jelenti, hogy a politikai, társadalmi, gazdasági, információs és katonai tényezők közötti kölcsönhatásokat is azonosítani kell. Például a gazdasági ellátási lánc sérülékenysége társadalmi elégedetlenséget okozhat, ami információs térben erősíti az ellenség narratíváját, és végül katonai mozgástér-szűküléshez vezet. A rendszerleírás célja nem az akadémiai teljesség, hanem a „beavatkozási pontok” azonosítása a későbbi design (COG/DC) számára.

A modern OE-elemzésben külön figyelmet érdemel az információs környezet (information environment). A COPD és kapcsolódó NATO gyakorlatok szerint az információs tér nem csak kommunikáció: tartalmazza az információ infrastruktúráját, a szereplők kommunikációs képességeit, a közönségek percepcióit és a befolyásolási műveletek logikáját. A tervezés során ezért célszerű külön alrészben rögzíteni: kulcs narratívák, fő csatornák, döntéshozók és közönségek, valamint a mérhetőség (mit tekintünk „változásnak” a percepcióban).

B.4. Szereplők, érdekeltségek és viszonyok

A szereplők (actors) listája a COPD-ben akkor értékes, ha az érdekeltségeket és a viszonyokat is rögzíti. A pusztá felsorolás helyett a törzsnek azt kell megértenie, ki mit akar, miben képes cselekedni, és milyen ösztönzők vagy korlátok befolyásolják. A viszonyok térképezése segít a branch/sequel opciókban: ha egy kulcsszereplő átáll vagy semleges marad, a COA-k kimenete megváltozhat.

- Jó gyakorlat: a szereplőknek „leverage” mező hozzáadása (mivel lehet rájuk hatni).
- Jó gyakorlat: a szereplők kapcsolatait nem csak barát–ellenség tengelyen ábrázolni, hanem függőségek és erőforrások mentén is.
- Tipikus hiba: a partner erők képességeinek túlértékelése, ami később tranzíciós kockázatot okoz.

B.5. Ellenség modell: COA-k, sebezhetőségek és adaptáció

A COPD a hírszerzési támogatást a tervezés integrált elemének tekinti. Az ellenség modellnek legalább két COA-t kell tartalmaznia: a legvalószínűbbet (most likely) és a legveszélyesebbet (most dangerous). A COA-k összevetése segít a risk appetite rögzítésében:

a parancsnoknak látnia kell, melyik ellenség viselkedés esetén milyen kockázatot vállal a választott megközelítés. Az ellenség modell minőségének egyik jó indikátora, hogy képes-e adaptációt leírni: mikor és hogyan változtat a taktikán, milyen jelzésekből lehet ezt felismerni, és melyik DC/indikátor érintett.

B.6. End state, success criteria, exit criteria

Az end state a művelet kívánt befejező állapota, amelyhez success criteria tartozik. A success criteria akkor használható, ha mérhető és időben értelmezhető: nem elég azt írni, hogy „stabilitás nő”, hanem azt kell rögzíteni, milyen konkrét, több forrásból ellenőrizhető jelek alapján tekinthető a stabilitás megfelelőnek. A COPD logikájában az exit criteria a success criteria operatív része: milyen feltételek mellett történik az átadás, mi minősül fenntartható állapotnak, és mi a minimum, ami nélkül a tranzíció nem indokolt.

Gyorsított tervezésben az end state gyakran „vázlatos” marad, de a COPD szerint még ilyenkor is szükséges egy minimális success criteria készlet, mert enélkül az assessment és a Phase 6 tranzíció nem tud működni. A success criteria-hoz célszerű felelősöket rendelni (ki méri), és előre rögzíteni, milyen adatokból állítjuk elő a képet (ISR, OSINT, partner jelentések). Ezzel csökkenthető a „visszatekintő igazolás” torzítása, amikor a művelet végén a törzs utólag próbálja igazolni, hogy a célok teljesültek.

B.7. SSA minőségbiztosítás: gyors checklist

- Problem statement–end state–success criteria logikai lánc konzisztens?
- Korlátozások és feltételezések listája teljes, owner és validációs terv van?
- Ellenség COA-k és adaptációs jelek szerepelnek, CCIR-hez kötve?
- OE rendszerleírásban megjelennek a kölcsönhatások, nem csak dimenzió-lista?
- SSA röviden is érthető: a parancsnok 10 perc alatt képet kap a döntési kérdésről?

Melléklet C – MRO/COA csomag: sablonok, pontozás és dokumentált bizonyíték

A MRO/COA csomag a parancsnoki döntés közvetlen inputja. A COPD szerint a COA-nak nem kell túl részletesnek lennie a Phase 3-ban, de elég strukturáltnak kell lennie ahhoz, hogy wargamingre és összehasonlításra alkalmas legyen. A melléklet két célt szolgál: egyrészt részletes COA-sablont ad (mezőkkel és kitöltési kérdésekkel), másrészt bemutatja, hogyan lehet a pontozást „bizonyítékhoz kötni”, hogy a COA-összevetés ne pusztán véleményösszegezés legyen.

C.1. COA leíró sablon (COA narrative template)

A COA-narratíva ideális esetben 1–2 oldal. A COPD-ben a túl hosszú COA leírás elveszíti a döntéstámogató funkcióját, mert a parancsnoki brief idejébe nem fér bele. A COA-narratíva mezői: (1) Approach (egy mondatban); (2) Main effort és a művelet súlypontja; (3) Phasing és critical events; (4) Decisive conditions, fő hatások és LoE/LoO hozzárendelés; (5) Erő- és képesség-igény; (6) LOG/CIS előfeltételek; (7) Civil és információs következmények; (8) Kockázatregiszter top tételek; (9) Assessment: 3–5 kulcsindikátor.

Mező	Kitöltés
Approach (1 mondat)	
Main effort	
Phasing + critical events	
Kulcs DC-k + LoE/LoO	
Erő/képesség igény	
LOG előfeltételek	
CIS/C2 előfeltételek	
Civil/Info hatások	
Top kockázatok + mitigáció	
Kulcs MOP/MOE (3–5)	

Melléklet C.1. COA narrative template (saját szerkesztés).

C.2. Pontozási fegyelem: súlyozás és bizonyíték

A COA comparison során a pontozás akkor segíti a döntést, ha a törzs előre rögzíti a súlyokat (a parancsnoki prioritások alapján), és minden pontszámhoz rövid bizonyítékot rendel. A bizonyíték lehet wargame esemény, logisztikai számítás, jogi elemzés vagy hírszerzési értékelés. A „bizonyíték mező” célja a vita minőségének javítása: ha nincs

bizonyíték, akkor a pontszám feltételezés, és ezt explicit módon kezelni kell (assumption register).

A pontozás tipikus kognitív csapdája a horgonyzás (anchoring): az elsőként bemutatott COA „alapértelmezett” lesz, és a többi ehhez képest kap pontszámot. A COPD gyakorlatban ezért több törzs rotálja a prezentációs sorrendet, vagy előre rögzíti, hogy a pontozás kritériumonként történik (azaz minden COA-nál ugyanazt a kritériumot vizsgáljuk, mielőtt továbblépünk). Ez csökkenti a sorrend hatását, és növeli az összehasonlíthatóságot.

C.3. Wargame dokumentálás: minimum outputok

A wargame dokumentálása gyakran alulértékelt, pedig a Phase 4 döntéstámogatás (DST/DSM) és a későbbi lessons learned is erre épül. A wargame minimum outputjai a COPD logikájában: (1) döntési pontok és triggererek; (2) CCIR és információgyűjtési prioritások; (3) feltárt gap-ek és képességhiányok; (4) javasolt mitigációk és branch/sequel opciók; (5) frissített kockázatregiszter. Ha ezek nincsenek rögzítve, a wargame a gyakorlatban „beszélgetés” marad, és nem lesz hatása a tervre.

- Jó gyakorlat: a wargame során minden critical eventhez külön sor a record sheeten, DP/CCIR mezőkkel.
- Jó gyakorlat: a red team és a blue team következtetéseinek elkülönítése (mi tény, mi vélemény).
- Tipikus hiba: a wargame csak a legvalószínűbb ellenség COA-ra készül, a legveszélyesebbre nem.

Melléklet D – CONOPS → OPLAN átmenet: termékek, annex-index és végrehajthatósági tesztek

A Phase 4 két alrésze közötti átmenet (CONOPS-ról OPLAN/OPORD-ra) a tervezés egyik legnagyobb „minőségi kockázata”. A CONOPS gyakran koherens narratíva, de az OPLAN akkor lesz végrehajtható, ha a szakági annexek ténylegesen összhangban vannak a base plannel. A melléklet olyan ellenőrző kérdéseket és egy annex-index sablont ad, amely segít végrehajthatósági tesztet futtatni: a terv logikailag, erőforrásban, C2-ben és jogilag „összeáll-e”.

D.1. Annex-index (példa struktúra)

Annex	Tartalom (rövid)	Kapcsolódó döntési pont/indikátor
A	Task organization, C2, command relationships	DP-k: C2 változások
B	Intelligence, ISR plan, PIR/CCIR	CCIR-k; trigger adatok
C	Operations: scheme of manoeuvre/fires	DC-k; critical events
D	Logistics & sustainment	Culmination; DP-LOG
E	Personnel	rotáció; veszteség
F	Public affairs/STRATCOM	IO indikátorok; reputációs DP
G	Civil-military cooperation (CIMIC)	civil hatás MOE-k
H	Legal/ROE	acceptability; targeting szabályok
K	Communications and Information Systems (CIS)	C2 redundancia; cyber DP

Melléklet D.1. Annex-index példa (saját szerkesztés).

D.2. Végrehajthatósági tesztek (feasibility tests)

A végrehajthatóság vizsgálata célszerűen négy dimenzióban történik. (1) Idő: a phasing és critical events ütemezése reális-e, figyelembe véve a telepítést, fenntartást és rotációt. (2) Erőforrás: a rendelkezésre álló erők és képességek elegendők-e a main effort támogatásához;

a hiányok kezeltek-e (mitigáció, partner erők, prioritások). (3) C2/CIS: a parancsnoki kapcsolatok világosak-e és a kommunikáció védett, interoperábilis-e; a degradált működés kezelése (fallback) rögzített-e. (4) Jogi/politikai: a ROE és mandátum ténylegesen lefedi-e a tervben szereplő tevékenységeket; a nemzeti caveat-ek kompatibilisek-e a feladat-hozzárendeléssel.

- Feasibility teszt kérdés: melyik annex „bukik meg” elsőként stressz alatt (LOG, CIS vagy Legal)?
- Feasibility teszt kérdés: ha 20–30% erő kiesik, a terv mely része omlik össze, és van-e branch?
- Feasibility teszt kérdés: ha az ellenség adaptációja miatt a fő indikátorok romlanak, milyen DP-k aktiválódnak?

Melléklet E – DST/DSM módszertan: döntési pontok, triggerek és CCIR összerendelése

A Decision Support Template/Matrix a COPD egyik legnagyobb gyakorlati értéke: előre strukturálja a döntést. A DST/DSM akkor működik, ha a döntési pontok nem „mindenre” vonatkoznak, hanem konkrét, parancsnoki szintű döntésekre. A triggereknek mérhetőnek kell lenniük, a CCIR pedig azt rögzíti, milyen információ nélkül nem hozható meg a döntés. A melléklet módszertani lépései a DST/DSM felépítéséhez: (1) wargame critical events; (2) döntési kérdések azonosítása; (3) DC-hez kötés; (4) indikátor és küszöb; (5) adatforrás és felelős; (6) opciók és ajánlás.

E.1. Döntési pont katalógus – „kevesebb, de jobb” elv

A döntési pont katalógus összeállításakor a leggyakoribb hiba a túl sok DP. Ha minden apró változás DP, a reporting rendszer túlterhelődik, és a parancsnok információs zajt kap. A COPD logika szerint DP csak akkor indokolt, ha (a) a döntés a parancsnok szintjén van, (b) a döntésnek több érdemi opciója van, és (c) a döntésnek jelentős kockázat- vagy erőforrás-következménye van. A DP-ket célszerű a design kulcselemeihez (DC-khez) kötni, mert így a DP-k száma természetes módon korlátozott és a traceability erős.

Melléklet F – Assessment mélyítés: adatforrások, confidence level és trianguláció

Az assessment minősége nem csak azon múlik, milyen indikátorokat választunk, hanem azon is, mennyire megbízhatóak az adatforrások és a következtetés. A Commander's Handbook hangsúlyozza a confidence level rögzítését: ugyanaz a trend más döntési következményt indokolhat, ha a bizonyosság alacsony. A melléklet célja, hogy gyakorlati módszert adjon a confidence és trianguláció kezelésére úgy, hogy a döntési briefekben ez gyorsan áttekinthető legyen.

F.1. Adatforrás-típusok és tipikus torzítások

ISR források erőssége a részletesség és időbeliség, de korlátozás a besorolás és a lefedettség. OSINT gyors és széles körű, de manipulálható és zajos. Partner jelentések helyi tudást adhatnak, de eltérő standardokkal és politikai szűrőkkel érkehetnek. Civil szervezetek adatbázisai értékesek lehetnek humán indikátorokhoz, de metodikai különbségek miatt validáció szükséges. A tervezés során célszerű minden kulcsindikátorhoz legalább két, lehetőleg eltérő típusú forrást rendelni, és rögzíteni, mi számít „megerősítésnek”.

F.2. Confidence level skála – javasolt egyszerűsítés

Gyakorlati környezetben a túl részletes bizonyossági skála használhatatlan. A bevált megközelítés egy háromfokozatú skála: High/Medium/Low. High akkor, ha több, egymástól független forrás megerősíti a trendet, és a torzítás kockázata alacsony. Medium, ha van megerősítés, de a források nem teljesen függetlenek vagy a torzítás kockázata közepes. Low, ha a trend egyetlen forrásból vagy nagyon zajos adatokból látható. A döntési briefben a confidence jelölése segít a parancsnoknak a kockázatvállalásban: alacsony confidence esetén gyakran további CCIR és gyűjtés indokolt, mielőtt nagy erőforrás-átrendezés történik.

F.3. Trianguláció a DC mérésében

A DC-k jellemzően összetett állapotok, ezért egyetlen MOE ritkán elég. A trianguláció azt jelenti, hogy különböző indikátorokból állítunk össze következtetést, és az indikátorokat nem egyformán súlyozzuk. Például egy „partner legitimáció erősödik” DC esetén a leading indikátor lehet a partner erők toborzási trendje és a civil együttműködés jelei, míg a lagging indikátor lehet a visszatérő incidensek csökkenése és a helyi bizalom mérése. A triangulációt segíti, ha a LoE owner felel a DC-hez tartozó indikátorcsomagért, és a reporting cycle-ben egyetlen integrált narratívát ad, nem pedig külön adatlistákat.

Melléklet G – Gyors referenciakártyák (Phase 1–6) – tervezési minimum

A melléklet „zseb-checklist” jelleggel, fázisonként összefoglalja a tervezési minimumot. Célja, hogy gyorsított tervezésben is biztosítható legyen a döntési transzparencia és a traceability. A listák nem teljesek, de a leggyakoribb hibák elkerülését támogatják.

G.1. Phase 1 (ISA) – minimum

- Trigger és döntési kérdés rögzítése (1 mondat).
- JOPG felállítása és battle rhythm alap (review/brief/döntés).
- Gap/CCIR top 5 + owner + határidő.
- Dokumentumkezelés: baseline mappa + verziókód + change log.
- Első risk register vázlat (Top 5).

G.2. Phase 2 (SSA) – minimum

- Problem statement, end state, success criteria (mérhető).
- Korlátozások/feltételezések explicit listája, validációval.
- Ellenség COA-k (most likely/most dangerous) és adaptációs jelek.
- OE rendszerleírás: fő kölcsönhatások és beavatkozási pontok.
- Initial operational approach (vázlat).

G.3. Phase 3 (MRO/COA) – minimum

- 2–3 eltérő COA (approach + main effort + phasing).
- COA-khoz LOG/CIS előfeltételek és kockázat top tételek.
- Wargame record sheet kitöltve (A-R-C) critical eventenként.
- DP-k és CCIR-ek azonosítva, első DST/DSM vázlat.
- COA comparison mátrix bizonyítékkal.

G.4. Phase 4 (CONOPS/OPLAN) – minimum

- CONOPS jóváhagyva: DC/LoE/LoO + indikátorok + phasing.
- OPLAN feladat-hozzárendelés és ütemezés (sync matrix).
- Annexek integráltan készülnek (LOG/CIS/Legal/Info).
- DST/DSM véglegesítve (triggerek + küszöbök + CCIR).
- Feasibility teszt lefutott (idő/erőforrás/C2/jog).

G.5. Phase 5 (Execution) – minimum

- Assessment battle rhythm-ben, trend + confidence jelöléssel.

- FRAGO változásnapló és traceability (miért, mire hat).
- Risk register és assumptions folyamatos frissítés.
- Döntési pontok monitorozása és CCIR státusz követése.
- LL megfigyelések rögzítése (observation log).

G.6. Phase 6 (Transition) – minimum

- Exit criteria mérhető és összhangban a success criteria-val.
- Átadási felelősségek, ütem, fenntarthatósági feltételek.
- Final assessment és post-operation review.
- LL action plan + validation (beépítés SOP/sablon/képzés).

Melléklet H – Cyber és információs dimenzió integrációja a COPD-tervezésben

A COPD v3.1 alaplogikája domain-agnosztikus: a design elemek (DC, LoE/LoO, DP, CCIR, assessment) bármely domainben értelmezhetők. Ugyanakkor a cyber és az információs tevékenységek integrációja gyakran külön módszertani kihívásokat hoz, mert (1) a hatások nehezen mérhetők, (2) a műveletek időskálája eltérhet a kinetikus műveletektől, (3) a jogi és politikai érzékenység magas, valamint (4) az attribúció és a bizonyosság gyakran alacsony. A melléklet azt mutatja be, hogyan lehet a COPD keretein belül ezeket a sajátosságokat kezelni anélkül, hogy a tervezés „külön világokra” szakadna.

H.1. Cyber hatások és decisive conditions – mérhetőség és proxyk

Cyber műveleteknél a „hatás” gyakran nem fizikai és nem azonnal látható. Például egy ellenség C2 hálózatának zavarása nem feltétlenül jelenik meg azonnal a csatatéren, és a megfigyelhető jelek (például rádiócsend, alternatív csatornák használata) csak proxy indikátorok. A COPD-hez illeszkedő megközelítés az, hogy a cyber tevékenységet nem önmagában, hanem a DC-hez rendelve tervezzük: mi az a döntési állapot (DC), amelyhez a cyber hozzájárul, és milyen mérhető jel (MOE) utal arra, hogy a DC felé mozdultunk. A mérést célszerű több szintre bontani: technikai MOP (végrehajtottuk-e), operációs MOP (az ellenség rendszerében történt-e változás), és műveleti MOE (a döntési ciklus lassult-e, a támadás/reakció ideje nőtt-e). Ez a rétegezés csökkenti azt a kockázatot, hogy a technikai sikerességet automatikusan stratégiai hatásnak tekintsük.

H.2. Információs környezet és STRATCOM – LoE owner és indikátor-fegyelem

Az információs tevékenységeknél a leggyakoribb hiba, hogy az IO/STRATCOM „mindenhez hozzászól”, de nincs egyértelmű felelősségi rend. A COPD szemléletében ezért célszerű az információs LoE-nek tulajdonost (LoE owner) adni, aki felel az indikátor-katalógusért, az adatminőségért és a reporting narratívájáért. A mérésnél a „van üzenetünk” nem indikátor; indikátor csak akkor értelmezhető, ha konkrét közönséghez és változáshoz kötődik. Gyakorlatias megoldás az indikátorok három kategóriája: reach (elérés), engagement (interakció) és perception/behaviour (percepció vagy viselkedés változása). A harmadik a legértékesebb, de a legnehezebb; ezért a COPD-logika szerint a reach/engagement inkább leading indikátor, míg a perception/behaviour inkább lagging.

H.3. Jog és politika: acceptability „korai beégetése”

Cyber és információs műveleteknél a jogi és politikai elfogadhatóság (acceptability) nem a terv végén kezelendő. A COPD keretében a Legal és STRATCOM inputnak már a Phase 2–

3 során meg kell jelennie, különben a COA-k összehasonlítása torz lesz: a törzs olyan COA-t preferálhat, amely technikailag hatékony, de politikailag vállalhatatlan. A jó gyakorlat a „korai vörös vonalak” rögzítése (red lines): milyen célpontok, milyen időzítés vagy milyen attribúciós kockázat esetén a COA elfogadhatatlan. Ezek a red line-ok később DP-kké alakíthatók (például reputációs küszöb vagy nemzeti caveat aktiválódása), így a végrehajtás során is transzparens marad, miért és mikor szükséges váltani.

H.4. CCIR cyber/IO kontextusban – „mit nem tudunk, ami nélkül nem dönthetünk”

A CCIR cyber/IO területen gyakran túl technikai, ezért a parancsnok számára nem döntésképes. A COPD szellemében célszerű a CCIR-t döntési kérdéshez kötni: milyen információ nélkül nem dönthet a parancsnok. Például nem az a CCIR, hogy „milyen sérülékenységi van a rendszerben”, hanem az, hogy „a műveleti időablakban várható-e az ellenség C2 redundancia aktiválása”, vagy „a reputációs kár eléri-e a politikai tűréshatárt”. A technikai részletek PIR-ként (priority intelligence requirement) megjelenhetnek, de a parancsnoki briefben a döntési relevancia a lényeg.

Melléklet I – Szemléltető tervezési vignetta: a COPD-termékek összekapcsolása egy rövid forgatókönyvben

A vignetta célja nem egy valós művelet bemutatása, hanem annak szemléltetése, hogyan kapcsolódnak egymáshoz a COPD-termékek (SSA → MRO/COA → CONOPS → OPLAN → DST/DSM → assessment) egy rövid, fiktív helyzetben. A vignetta különösen a traceability gondolkodást erősíti: minden design elemhez (DC, LoE) kapcsolódik feladatcsomag és mérőszám, és a döntési pontok előre rögzítettek.

1.1. Kiinduló helyzet (ISA/SSA vázlat)

Forgatókönyv: egy szomszédos államban fegyveres csoportok destabilizálják a határ menti régiót, amely humanitárius válsággal, menekültáramlással és információs műveletekkel párosul. A NATO politikai célja a regionális stabilitás helyreállítása és a menekültáramlás csökkentése, miközben a mandátum korlátozott: a művelet csak a határ menti biztonsági zónára terjed ki, és a civil infrastruktúra védelme kiemelt constraint.

SSA end state (vázlat): a biztonsági zóna ellenőrzött, a humanitárius segély eljut a célterületre, és a helyi kormányzati struktúrák képesek fenntartani a rendet. Success criteria (példák): (1) a fő utánpótlási útvonalakon az incidensek száma tartósan csökken; (2) a humanitárius konvojok rendszeresen és veszteség nélkül közlekednek; (3) a helyi biztonsági erők önállóan végrehajtanak járőrözési és checkpoint feladatokat; (4) a dezinformációs kampányok hatása mérhetően csökken (percepció MOE).

1.2. Design vázlat: COG, DC-k, LoE-k

A design során a törzs az ellenség COG-jaként azonosítja a fegyveres csoport „mobil erőkoncentrációs képességét” (gyors átcsoportosítás és rajtaütések). A CC lehet a gyors mozgás és rejtett logisztikai hálózat, a CR lehet a kulcsútvonalak és támogatói infrastruktúra, a CV pedig a szűk keresztmetszetek (hidak, üzemanyag-ellátás) és a kommunikációs függőségek. A DC-k például: DC1 – a fő útvonalak biztonsága; DC2 – az ellenség logisztikai utánpótlás korlátozása; DC3 – a helyi erők képességének növelése; DC4 – az információs térben a NATO/partner narratíva hitelességének erősítése. A LoE-k (security, sustainment, governance/partnering, information) ezekhez a DC-khez rendelhetők.

1.3. COA-k röviden és döntési logika

COA-A (erőkoncentrált): gyors telepítés, erős kinetikus jelenlét a fő útvonalak mentén, intenzív ISR + fires, célzott csapások a logisztikai hálózatra. Előny: gyors security hatás, erős

deterrence. Kockázat: civil károk és reputációs kár, magas LOG igény. COA-B (partner-központú): a NATO főleg training/advising, partner erők vezetik a műveleteket, NATO ISR és gyorsreagálás támogat. Előny: politikailag elfogadhatóbb, tranzíció könnyebb. Kockázat: lassabb security hatás, partner megbízhatóság. COA-C (information+denial): hangsúly a mobilitás megtagadásán (útvonalak kontrollja, akadályok), cyber/IO kombináció az ellenség C2 ellen, korlátozott kinetikus. Előny: alacsonyabb kinetikus kockázat; hátrány: cyber/IO bizonytalanság, mérhetőség.

1.4. DST/DSM példa: reputációs DP és LOG DP

A wargaming során két parancsnoki döntési pont emelkedik ki. DP-IO: reputációs incidens küszöb felett – kérdés, szükséges-e ROE módosítás és IO kampányváltás. Trigger: civil incidensek száma és OSINT trend (confidence jelöléssel). Opciók: (1) műveleti tempo csökkentése, (2) célkijelölési szabályok szigorítása, (3) proaktív STRATCOM + vizsgálat. DP-LOG: készletszint és útvonal-biztonság romlása – kérdés, szükséges-e tempo módosítás vagy alternatív útvonal. Trigger: készletszint küszöb alatt + útvonalon jelentett incidensek. Ezek a DP-k a Phase 5-ben gyorsítják a döntést, mert az opciók és kockázatok előre elemzettek.

1.5. Assessment vázlat és tranzíció

Az assessmentben a security LoE-hez MOE lehet az incidens-trend és a konvojok sikerességi aránya, governance/partnering LoE-hez a partner erők önálló műveleti aránya és a közös műveletek sikeressége, information LoE-hez pedig a bizalom/percepció mérése és a dezinformációs tartalmak terjedésének csökkenése. A tranzíció (Phase 6) akkor indítható, ha a success criteria trendje stabil: nem egyszeri jó érték, hanem tartós teljesülés. A vignetta üzenete: a COPD-termékek akkor adnak valódi értéket, ha a design elemeihez (DC, LoE) végig hozzá vannak rendelve a feladatok, a döntési pontok és az indikátorok.

RÖVIDÍTÉSEK JEGYZÉKE

1. Az alábbi rövidítéstáblázat a COPD v3.1-ben használt rövidítéseket tartalmazza. Ahol egy rövidítést széles körben használnak a NATO-szerte, ott a COPD-vel kapcsolatos fő forrásdokumentumot idéztük forrásként.
2. A hivatalos NATO-doktrínákban használt rövidítéseket és terminológiát központilag, a NATO központjában, a NATO Szabványügyi Hivatal NATO Terminológiai Irodája kezeli.

RÖVIDÍTÉS	JELENTÉS	FORRÁS
-----------	----------	--------

A

AAP	Allied Administrative Publication	AAP 15
AB	Assessment Board	
ACO	Allied Command Operations	AAP 15
ACOS	Assistant Chief of Staff	AAP 15
ACT	Allied Command Transformation	AAP 15
ACTORD	Activation Order	AAP 15
ACTPRED	Activation Pre-deployment	NCRSM
ACTREQ	Activation Request	AAP 15
ACTWARN	Activation Warning	AAP 15
ADAMS	Allied Deployment and Movement System	AAP 15
AD	ACO Directive	
ADL	Allied Disposition List	AAP 15
ADM	Accelerated Decision Making	NCRSM
AFL	Allied Forces List	
AIFS	Allied Information Flow System	AAP 15
AIMS	AIFS Integrated Message System	AAP 15
AIRCOM	(Headquarters) Allied Air Command	MC 0324
AJD	Allied Joint Doctrine	AAP 15

RÖVIDÍTÉS	JELENTÉS	FORRÁS
AJP	Allied Joint Publication	AAP 15
AMCC	Allied Movement Coordination Centre	AAP 15
AOI	Area of Interest	AAP 15
AOM	Allied Operations and Missions	
AOO	Area of Operations	AAP 15
AOR	Area of Responsibility	AAP 15
APOD	Airport of Debarkation	AAP 15
APOE	Airport of Embarkation	AAP 15
ARRP	AOM Requirements and Resources Plan	
ASG	Assistant Secretary General	AAP 15
AU	African Union	
AWG	Assessment Working Group	
B		
BICES	Battlefield Information Collection and Exploitation System	AAP 15
BI	Building Integrity	
Bi-SC	Bi-Strategic Commands	AAP 15
BMD	Ballistic Missile Defence	
C		
C2	Command and Control	AAP 15
C2S	Command and Control System	AAP 15
C2W	Command and Control Warfare	AAP 15
CA	Comprehensive Approach	COPD
CAAC	Children and Armed Conflict	MCM 0016-2012
CALM	Crane Attachment Lorry Mounted	
CAPAC	Consultation, Approval, Promulgation and Activation Procedures	
CAT	Cross Functional Action Team (Plans or Ops) Replaced by Strategic Operations Planning Group	
CAT	Campaign Assessment Tool	TOPFAS
CBRN	Chemical, Biological, Radiological and Nuclear	AAP 15

RÖVIDÍTÉS	JELENTÉS	FORRÁS
CC	Component Command	AAP 15
CC	Critical Capability	
CCD	CIS and Cyber Defence	
CCIR	Commander's Critical Information Requirement	AAP 15
CCIRM	Collection, Coordination and Intelligence Requirements Management	AAP 15
CD	Cyber Defence	
CD	Collective Defence	
CE	Crisis Establishment	AAP 15
CEP	Civil Emergency Planning	AAP 15
CEPC	Civil Emergency Planning Committee	AAP 15
CG	Command Group	
CIMIC	Civil-Military Cooperation	AAP 15
CIS	Communications and Information Systems	AAP 15
CIVAD	Civil Actors Advisor	
CIVCAS	Civilian Casualties	
CJSOR	Combined Joint Statement of Requirements	AAP 15
CLE	CyOC Liaison Element	
CM	Crisis Management	
CMALT	Civil-Military Assessment and Liaison Team	
CMC	Chairperson of the Military Committee	AAP 15
CMI	Civil-Military Interaction	
CMPS	Civil-Military Planning and Support Section	
CMTF	Crisis Management Task Force	
CMRB	Crisis Management Requirements Board	
CN	Contributing Nation	AAP 15
CNA	Computer Network Attack	AAP 15
CND	Computer Network Defence	
CNI	Critical National Infrastructure	
CNMA	Complementary Non-Military Activity	COPD

RÖVIDÍTÉS	JELENTÉS	FORRÁS
CoA	Course of Action	AAP 15
CoC	Chain of Command	
CoG	Centre of Gravity	AAP 15
COM	Commander	AAP 15
COMPASS	Comprehensive Approach Specialist Support	
CONOPS	Concept of Operations	AAP 15
CONPLAN	Contingency Plan	AAP 15
COP	Contingency Operations Plan	CCOM HB
COPD	Comprehensive Operations Planning Directive	MC 0133
CORSOM	Coalition Reception Staging and Onward Movement	LOGFAS
COS	Chief of Staff	AAP 15
CPoE	Comprehensive Preparation of Operational Environment	AJP 05
CPP	Cultural Property Protection	
CR	Critical Requirements	
CRD	Commander's Required Date	AAP 15
CRM	Crisis Response Measure	AAP 15
CRO	Crisis Response Operation	AAP 15
CRP	Crisis Response Planning	
CSAR	Combat Search and Rescue	AAP 15
CSC	Convoy Support Centre	
CSO	Contractor Support to Operations	
CT	Counter Terrorism	
CTS	COSMIC Top Secret	AAP 15
CUoE	Comprehensive Understanding of the Environment	
CV	Critical Vulnerabilities	
CyOC	Cyber Operations Centre	
D		
DAMCON	Damage Control	AAP 15
DB	Decision Brief	

RÖVIDÍTÉS	JELENTÉS	FORRÁS
DC	Decisive Condition	AJP 01
DCIS	Deployable Communication and Information Systems	AAP 15
DCOS	Deputy Chief of Staff	AAP 15
DCP	Data Collection Plan	
DDP	Detailed Deployment Plan	AAP 15
DDR	Disarmament, Demobilisation and Reintegration	AAP 15
DE	Deployable Element	
DHS	Document Handling System	
DIME	Diplomatic, Information, Military, Economic	
DIMEFIL	Diplomatic, Information, Military, Economic, Financial, Intelligence, Legal	
DIRLAUTH	Direct Liaison Authority	
DoA	Desired Order of Arrival	AAP 15
DP	Decision Point	
DPRE	Displaced Persons and Refugees	
DSACEUR	Deputy SACEUR	AAP 15
DSM	Decision Support Matrix	
DTG	Date-Time Group	AAP 15
E		
EADRCC	Euro-Atlantic Disaster Response Coordination Centre	AAP 15
EAPC	Euro-Atlantic Partnership Council	AAP 15
EEFI	Essential Elements of Friendly Information	AAP 15
EOD	Explosive Ordnance Disposal	AAP 15
EOR	Explosive Ordnance Reconnaissance	AAP 15
ERS	Enablement and Resilience Section	
E&T	Exercises & Training	
EU	European Union	AAP 15
EUMS	European Union Military Staff	AAP 15
EVAL	Evaluation	
EVE	Effective Viable Execution	LOGFAS

RÖVIDÍTÉS	JELENTÉS	FORRÁS
EW	Electronic Warfare	AAP 15
F		
FACCES	Feasibility, Acceptability, Completeness, Compliance, Exclusivity and Suitability	COPD
FAD	Force Activation Directive	MC 0133
FC	Functional Commanders	
FCE	Forward Coordination Element	CFAO
FFIR	Friendly Forces Information Requirement	AAP 15
FGen	Force Generation	
FGC	Force Generation Conference	
FGR	Force Generation and Readiness Branch	
FMovPC	Final Movement Planning Conference	
FORCE PREP	Force Preparation	AAP 15
FP	Force Protection	AAP 15
FPG	Functional Planning Guide	AAP 15
FRAGO	Fragmentation Order	AAP 15
FS	Functional Services	
FTDM	Fast Track Decision-Making	MC 0133
G		
GCOP	Generic Contingency Operations Plan	
GENAD	Gender Advisor	
GO	Governmental Organisation	
GRF	Graduated Readiness Forces	
GRP	Graduated Response Plan	
H		
HA	Humanitarian Assistance / Aid	
HF	High Frequency	
HISAD	Historical Advisor	
HN	Host Nation	AAP 15
HNS	Host-Nation Support	AAP 15
HOTO	Handover Takeover	

RÖVIDÍTÉS	JELENTÉS	FORRÁS
HQ	Headquarters	AAP 15
HQ JFC	Headquarters Joint Force Command	MC 0324
HUMINT	Human Intelligence	AAP 15
HVA/A	High Value Asset / Area	
HVT	High Value Target	AAP 15
I		
I&IS	Information & Intelligence Sharing	
IAMD	Integrated Air and Missile Defence	
IC	International Community	
ICC	International Criminal Court	
ICE	Initial Command Element	CFAO
ICI	Istanbul Cooperation Initiative	AAP 15
ICRC	International Committee of the Red Cross	
IDB	Integrated Database	AAP 15
IE	Information Environment	
IEA	Information Environment Activities	
IEG	Information Exchange Gateway	MC 0593
IER	Information Exchange Requirement	AAP 15
IHL	International Humanitarian Law	
IKM	Information Knowledge Management	
IM	Information Management	
IMINT	Imagery Intelligence	AAP 15
IMovPC	Initial Movement Planning Conference	
IMS	International Military Staff	AAP 15
Info Ops	Information Operations	AAP 15
INTEL FS	Intelligence Functional Services	
INTREP	Intelligence Report	AAP 15
INTSUM	Intelligence Summary	AAP 15
IMINT	Imagery Intelligence	AAP 15
IMS	International Military Staff	AAP 15
IO	International Organisation	AAP 15

RÖVIDÍTÉS	JELENTÉS	FORRÁS
IoP	Instrument of Power	
IS	International Staff	AAP 15
ISB	Intermediate Staging Base	AAP 15
ISR	Intelligence, Surveillance and Reconnaissance	AAP 15
ISTAR	Intelligence Surveillance Target Acquisition and Reconnaissance	AAP 15
I&W	Indications and Warnings	
J		
JALLC	Joint Analysis and Lessons Learned Centre	AAP 15
JCB	Joint Coordination Board	
JCO	Joint Coordination Order	
JCOP	Joint Common Operational Picture	AD 80-80
JEWCS	Joint Electronic Warfare Core Staff	AAP 15
JFC	Joint Force Command	AAP 15
JHQ	Joint Headquarters	
JIA	Joint Implementation Arrangement	
JIPOE	Joint Intelligence Preparation of the Operational Environment	
JISD	Joint Intelligence and Security Division	
JLSG	Joint Logistics Support Group	
JOA	Joint Operations Area	AAP 15
JOC	Joint Operations Centre	AAP 15
JOPG	Joint Operations Planning Group	AAP 15
JPDAL	Joint Prioritised Defended Asset List	
JPTL	Joint Prioritised Target List	
JSEC	Joint Support Enabling Command	
JTF	Joint Task Force	
JTF HQ	Joint Task Force Headquarters	CFAO
JWC	Joint Warfare Centre	
K		
KB	Knowledge Base	
KD	Knowledge Development	

RÖVIDÍTÉS	JELENTÉS	FORRÁS
KI	Key Infrastructure	
KM	Knowledge Management	
KR	Knowledge Requirement	
L		
LANDCOM	Headquarters Allied Land Command	MC0324
LEGAD	Legal Advisor	AAP 15
LL	Lessons Learned	AAP 15
LN	Lead Nation	AAP 15
LO	Liaison Officer	
LoA	Level of Ambition	
LOAC	Law of Armed Conflict	
LoC	Lines of Communication	
LoE	Line of Effort	
LOGFAS	Logistic Functional Area Services	AAP 15
LoO	Line of Operations	AAP 15
M		
M&T	Movements and Transport	AAP 15
MAB	Mission Analysis Briefing	
MARCOM	Headquarters Allied Maritime Command	MC 0324
MB	Military Budget	AAP 15
MC	Military Committee	AAP 15
MD	Mediterranean Dialogue	AAP 15
MEDAD	Medical Advisor	
METOC	Meteorological and Oceanographic	AAP 15
MILENG	Military Engineering	AAP 15
MilPA	Military Public Affairs	
MJO	Major Joint Operation	
MMovPC	Main Movement Planning Conference	
MN	Multi National	
MN DDP	Multinational Detailed Deployment Plan	AAP 15
MOE	Measure of Effectiveness	AAP 15

RÖVIDÍTÉS	JELENTÉS	FORRÁS
MOP	Measure of Performance	AAP 15
MOU	Memorandum of Understanding	AAP 15
MOVCON	Movement Control	
MP	Military Police	AAP 15
MPC	Movement Planning Conferences	
MPEC	Military Political Economic Civil (sometimes written as PMEC)	COPD
MRO	Military Response Option	AAP 15
MSA	Military Strategic Action	COPD
MSE	Military Strategic Effect	COPD
MSO	Military Strategic Objective	COPD
MSR	Main Supply Route	
M&T	Movement and Transportation	
MVI	Mission Vital Infrastructure	
N		
NA5CRO	Non-Article 5 Crisis Response Operation	
NAC	North Atlantic Council	AAP 15
NAEW&CS	NATO Airborne Early Warning and Control System	AAP 15
NAGS	NATO Ground Surveillance System	
NATO HQ	NATO Headquarters	AAP 15
NC	NATO Confidential	AAP 15
NCIA	NATO Communication and Information Agency	
NCIRC	NATO Computer Incident Response Capability	
NCISG	NATO Communications and Information Systems Group	
NCRP	NATO Crisis Response Process	AAP 15
NCRS	NATO Crisis Response System	AAP 15
NCRSM	NATO Crisis Response System Manual	AAP 15
NCRTA	NATO Crisis Response Tracking Application	NCRSM
NCS	NATO Command Structure	AAP 15
NDPP	NATO Defence Planning Process	AAP 15

RÖVIDÍTÉS	JELENTÉS	FORRÁS
NED	NAC Execution Directive	NCRSM
NLLP	NATO Lessons Learned Portal	
NLLDB	NATO Lessons Learned Database	AAP 15
NED	NAC Execution Directive	NCRSM
NFS	NATO Force Structure	AAP 15
NFS JHQ	NATO Force Structure Joint Headquarters	CFAO
NGCS	NATO General Communications Systems	
NGO	Non-Governmental Organisation	AAP 15
NID	NAC Initiating Directive	NCRSM
NIFC	NATO Intelligence Fusion Centre	
NIMP	NATO Information Management Policy	
NIWS	NATO Intelligence Warning System	AAP 15
NMA	NATO Military Authority	AAP 15
NMIC	NATO Media Information Centre	
NMR	National Military Representative	AAP 15
NMSE	Non-Military Strategic Effect	
NMSO	Non-Military Strategic Objective	
NNTCN	Non-NATO Troop Contributing Nation	AAP 15
NOAH	NATO Operations Assessment Handbook	
NR	NATO Restricted	AAP 15
NRF	NATO Response Force	AAP 15
NS	NATO Secret	AAP 15
NSHQ	NATO Special Operations Headquarters	CFAO
NSIP	NATO Security Investment Programme	AAP 15
NSPA	NATO Support and Procurement Agency	
NU	NATO Unclassified	AAP 15
O		
OA	Operational Action	COPD
OCC E&F	Operational Capabilities Concept Evaluation and Feedback Programme	
OE	Operational Effect	COPD

RÖVIDÍTÉS	JELENTÉS	FORRÁS
OLRT	Operational Liaison and Reconnaissance Team	AAP 15
OO	Operational Objective	COPD
OPSA	Operations Assessment	
OPC	Operations Planning Committee	
OPCOM	Operational Command	AAP 15
OPCON	Operational Control	AAP 15
OPD	Operational Planning Directive	
OPC	Operations Policy Committee	AAP 15
OPG	Operational Planning Guidance	
OPLAN	Operations Plan	AAP 15
OPLE	Operational Planning and Liaison Element	
OPP	Operations Planning Process	AAP 15
OPORD	Operation Order	AAP 15
OPSA	Operations Assessment	
OPT	Operations Planning Tool	TOPFAS
ORBAT	Order of Battle	AAP 15
ORBATTOA	Order of Battle Transfer of Authority	
OSCE	Organisation for Security and Cooperation in Europe	AAP 15
OSO	Office of Special Operations	
OTF	Operational Task Force	
P		
PA	Public Affairs	AAP 15
PAO	Public Affairs Office	AAP 15
PASCAD	Public Affairs and Strategic Communications Advisor	
PASP	Political Affairs and Security Policy	
PCB	Police Capacity Building	
PCS	Preconditions for Success	COPD
PCWL	Potential Crisis Warning List	
PD	Partnership Directorate	
PD	Public Diplomacy	

RÖVIDÍTÉS	JELENTÉS	FORRÁS
PDD	Public Diplomacy Division	
PDIM	Primary Directive on Information Management	
PE	Peacetime Establishment	AAP 15
PfP	Partnership for Peace	AAP 15
PIR	Priority Intelligence Requirement	AAP 15
PM	Provost Marshal	AAP 15
PME	Political Military Estimate	NCRSM
PMEC	Political Military Economic Civil	
PMESII	Political Military Economic Social Infrastructure Information (i.e. Systems within the Engagement Space)	
PMR	Periodic Mission Review	NCRSM
PNS	Plan Numbering System	COPD
PO	Private Office	
POC	Point of Contact	AAP 15
PoC	Protection of Civilians	
POD	Port / Point of Debarkation	AAP 15
POE	Port / Point of Embarkation	AAP 15
POL	Petroleum, Oils and Lubricants	AAP 15
POLAD	Political Advisor	AAP 15
PoP	Point of Presence	AAP 15
PPI	Political Policy Indicators	
PPS	Political Policy Statement	
PsyOps	Psychological Operations	AAP 15
PVO	Private Volunteer Organisation	AAP 15
Q		
R		
RCB	Resources Coordination Board	
REM	Resource Management	
RES	Resources Directorate	
RFI	Request for Information	

RÖVIDÍTÉS	JELENTÉS	FORRÁS
RMP	Risk Management Process	
ROC	Rehearsal of Concept	
ROE	Rules of Engagement	AAP 15
ROEREQ	Rule-of-engagement Request	AAP 15
ROEIMP	Rules of Engagement Implementation Message	AAP 15
RoL	Rule of Law	
RoM	Rough Order of Magnitude	
ROTA	Release Other Than Attack	
RPOD	Rail Point of Debarkation	
RPOE	Rail Point of Embarkation	
RPPB	Resource, Planning and Policy Board	
RSOM (I)	Reception, Staging, Onward Movement (Integration)	AAP 15
RTR	Royal Tank Regiment	
S		
SA	Situational Awareness	
SAB	Situational Awareness Briefing	
SACEUR	Supreme Allied Commander Europe	AAP 15
SAC	Strategic Analysis Capability	
SACT	Supreme Allied Commander Transformation	AAP 15
SAR	Search and Rescue	AAP 15
SASE	Safe And Secure Environment	
SAT	Strategic Assessment Team	NCRSM
SAT	Systems Analysis Tool	TOPFAS
SATCOM	Satellite Communications	AAP 15
SC	Strategic Command	AAP 15
SCEPVA	Sovereign Cyber Effects Provided Voluntarily by Allies	
SCO	Strategic Coordination Order	COPD
SCPB	Strategic Communication Policy Board	
SCR	Senior Civilian Representative	AAP 15
SDP	Standing Defence Plan	AAP 15

RÖVIDÍTÉS	JELENTÉS	FORRÁS
SECGEN	Secretary General	AAP 15
SERP	Sequenced Response Plan	
SHAPE	Supreme Headquarters Allied Powers Europe	AAP 15
SHF	Secure High Frequency	
SIA	Strategic and International Affairs Advisor	
SIGINT	Signals Intelligence	AAP 15
SITCEN	Situation Centre	AAP 15
SJLSG	Standing Joint Logistic Support Group	
SLOC	Sea Lines of Communication	AAP 15
SMA	Strategic Military Advice	NCRSM
SMAP	Standard Manpower Procedure	MC0216
SME	Subject Matter Expert	
SN	Sending Nation	AAP 15
SO	Strategic Objective	
SOF	Special Operations Forces	
SOFA	Status of Forces Agreement	AAP 15
SOFAD	Special Operations Forces Advisor	
SOI	Standard Operating Instruction	
SOP	Standard Operating Procedure	AAP 15
SOPG	Strategic Operations Planning Group	
SOR	Statement of Requirements	AAP 15
SOSA	Systems of Systems Analysis	
SP	Stability Policing	
SPD	Strategic Planning Directive	COPD
SPMP	Strategic Political-Military Plan	NCRSM
SPOD	Seaport of Debarkation	AAP 15
SPOE	Seaport of Embarkation	AAP 15
SSA	SACEUR's Strategic Assessment	NCRSM
SSC	Single Service Command	MC 0324
SSI	Supported Supporting Inter-relationship	
STANAG	NATO Standardisation Agreement	AAP 15

RÖVIDÍTÉS	JELENTÉS	FORRÁS
StratCom	Strategic Communications	
SUPPLAN	Support Plan	AAP 15
T		
TA	Technical Agreement / Arrangement	
TASKORG	Task Organisation	
TBM	Theatre Ballistic Missile	AAP 15
TBMD	Theatre Ballistic Missile Defence	AAP 15
TCC	Theatre Component Command	
TCN	Troop Contributing Nation	AAP 15
TCSOR	Theatre Capability Statement of Requirements	AAP 15
TD	Targeting Directive	
TF	Task Force	
TG	Task Group	
TIM	Toxic Industrial Material	
TLB	Theatre Logistic Base	
TMD	Theatre Missile Defence	AAP 15
ToA	Transfer of Authority	AAP 15
TOO	Theatre of Operations	AAP 15
TOPFAS	Tools for Operations Planning Functional Area Services	AAP 15
ToR	Terms of Reference	AAP 15
TST	Time Sensitive Target	AD 80-70
U		
UCM	Urgent Capabilities Management	
UHF	Ultra-High Frequency	
UN	United Nations	AAP 15
UNOCHA	United Nations Office for the Coordination of Humanitarian Affairs	
UNSC	United Nations Security Council	
UNSCR	United Nations Security Council Resolution	
V		
VCOS	Vice Chief of Staff	

RÖVIDÍTÉS	JELENTÉS	FORRÁS
VI	Vital Infrastructure	
VTC	Video Teleconference	AAP 15
W		
WAN	Wide-Area Network	AAP 15
WISE	Web Information Services Environment	
WMD	Weapon of Mass Destruction	AAP 15
WNGO	Warning Order	
WP	Warning Problem	
X-Y-Z		