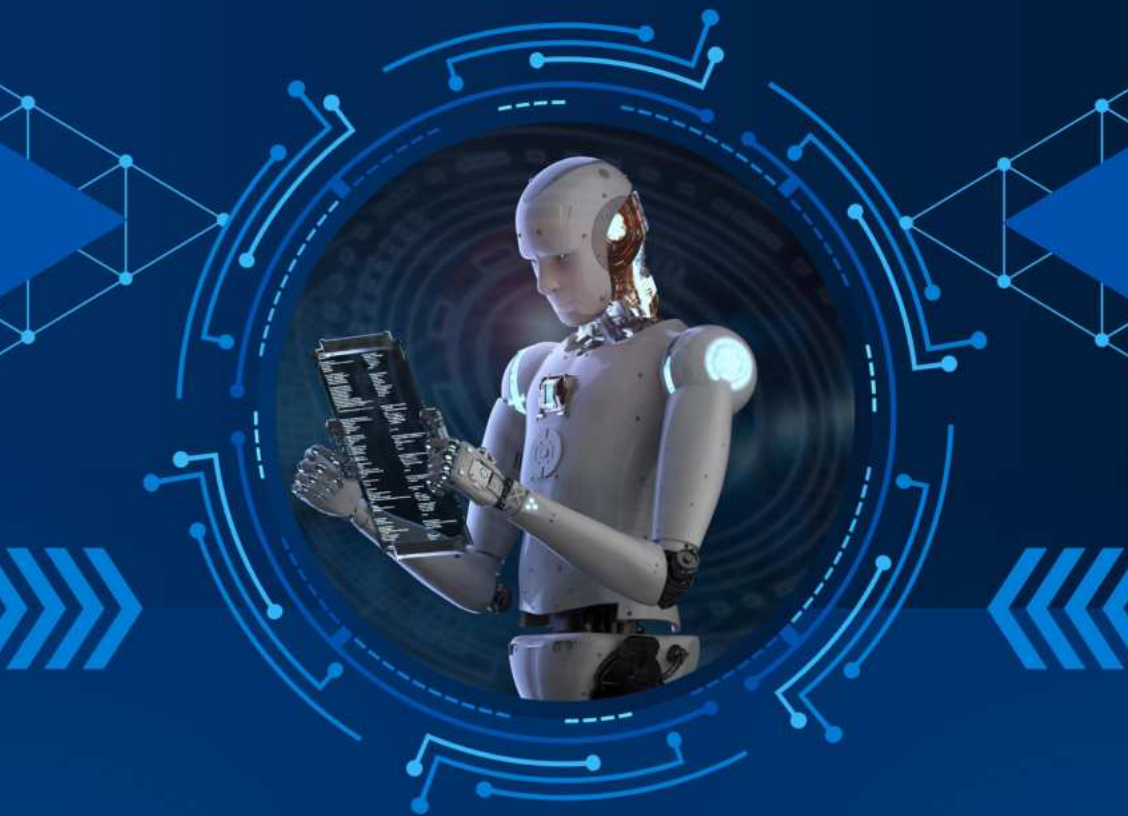




WHITE PAPER

INTELIGENȚA ARTIFICIALĂ ÎN SECURITATEA FIZICĂ

5 CAPABILITĂȚI FUNDAMENTALE

Ion IORDACHE



Ion IORDACHE, Ec.

Sunt economist licențiat în finanțe/bănci de profesie și absolvent al unei școli militare de ofițeri, specializat și certificat ca trainer profesionist și consultant de securitate. În prezent, îmi desfășor activitatea în cadrul companiei **RQM Certification** ca **Training and Development Manager**.

Am acumulat peste 30 de ani de experiență în securitatea privată, evaluarea riscurilor, managementul operațiunilor de securitate, securitatea informației și protecția datelor cu caracter personal, incluzând expertiză în **CPTED** (Crime Prevention Through Environmental Design) și educația adulților, fiind inițiatorul și autorul a două standarde ocupaționale, **“Consultant de securitate”** și **“Manager de securitate”**.

Dincolo de expertiza în securitatea fizică și informațională, sunt profund conectat la provocările viitorului. Investesc timp și rigoare analitică în cercetarea detaliată a **Inteligenței Artificiale**, fiind preocupat în mod special de evoluția acestui domeniu și de aspectele sale etice. Cred cu tărie că înțelegerea profundă a tehnologiei este vitală pentru a rămâne relevanți și responsabili într-o lume în continuă schimbare.



www.ioniordache.com



ion@ioniordache.com



www.rqmcert.ro



ion.iordache@rqmcert.ro

Februarie 2026 – V.0

Acest ghid a fost elaborat printr-o colaborare hibridă om-inteligență artificială, având la bază surse verificate, accesibile prin hyperlink-urile prezentate în Bibliografie. În procesul de documentare și creație au fost utilizate următoarele instrumente AI:

- ✓ **Perplexity AI**: pentru sinteza documentară și verificarea preliminară a surselor bibliografice.
- ✓ **Gemini Nano Banana**: pentru generarea imaginilor.
- ✓ **NotebookLm**: pentru generarea infograficelor.

Deși procesul a fost asistat tehnologic, responsabilitatea pentru selecția informațiilor, validarea veridicității acestora, precum și pentru opiniile și analizele exprimate îmi aparține în totalitate. Această notă subliniază angajamentul meu pentru o utilizare etică, transparentă și asumată a tehnologiei în cercetare.

CUPRINS

Introdúcere	04
Agentic AI	06
AI Agents	15
Generative AI	25
Neural Networks	35
Machine Learning	46
Cursuri de Specializare	56
Bibliografie	57

INTRODUCERE

Industria securității fizice trece prin cea mai profundă transformare din istoria sa, de la apariția camerei video până astăzi, o schimbare care depășește simpla supraveghere vizuală și se extinde spre capacitatea sistemelor de a înțelege contextul, de a anticipa riscurile și de a interveni în mod autonom.

Acest ghid traduce versatilitatea tehnologiei AI (Artificial Intelligence) în **5 piloni operaționali** concreți pentru securitatea fizică:

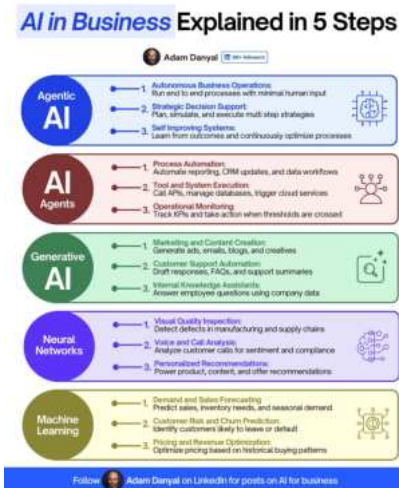
1. **Agentic AI (Creierul):** Sisteme care iau decizii strategice autonome.
2. **AI Agents (Mâinile):** Automatizarea execuției și integrarea sistemelor disparate.
3. **Generative AI (Vocea):** Transformarea procedurilor statice în asistenți conversaționali și training interactiv.
4. **Neural Networks (Simțurile):** Capacitatea de a vedea și auzi pericolele cu o precizie supraumană.
5. **Machine Learning (Predictția):** Trecerea de la securitate reactivă la securitate predictivă.

Materialul de față se adresează profesioniștilor din securitatea fizică, consultanți și manageri de securitate, evaluatori de risc, proiectanți și ingineri, și își propune să le ofere un traseu structurat și aplicabil pentru integrarea tehnologiei AI în activitatea lor, cu accent permanent pe eficiență operațională și pe gestionarea responsabilă a riscurilor asociate.

În același timp, ghidul este conceput pentru a fi accesibil și decidenților care beneficiază de serviciile acestor specialiști, fie că este vorba de proprietari care doresc să protejeze bunuri și persoane, fie de manageri și directori din organizații publice sau private care trebuie să evalueze investiții în securitate fizică bazată pe AI. Pentru aceștia, materialul demistifică tehnologia și oferă criterii concrete de evaluare a soluțiilor propuse de furnizori.

În peisajul actual al securității fizice, termenul "Inteligență Artificială" este invocat atât de frecvent și de necritic, încât a ajuns mai degrabă un instrument de promovare comercială decât o descriere fidelă a unor capacități tehnologice reale și verificabile. Ca specialist cu activitate dedicată educației și consultanței în domeniul securității, am ales să trec dincolo de discursul de marketing și să ofer o perspectivă tehnică riguroasă, practică și lipsită de ambiguitate.

Acest **White Paper** pornește de la cadrul conceptual excelent dezvoltat de **Adam Danyal** în "**AI in Business Explained in 5 Steps**". Mulțumesc autorului pentru generozitatea de a face publice aceste informații structurate.



Secțiunile care urmează transpun aceste categorii fundamentale din mediul de business în realitatea concretă a securității fizice, un domeniu guvernat de cerințe operaționale și constrângeri tehnice specifice. Cadrul teoretic a fost adaptat și reformulat în termeni cu care lucrează zilnic profesioniștii din securitate: evaluări de risc cu aplicabilitate practică, protocoale de intervenție pentru dispecerat, proceduri de pază cu criterii clare de verificare și arhitecturi de integrare a sistemelor de securitate.

IA în Afaceri: 5 Categoriile Cheie pentru Eficiență



AGENTIC AI

Agentic AI în Securitatea Fizică De la "Instrument" la "Partener"

În majoritatea cazurilor, când industria vorbește despre "**AI în securitate**", se referă la algoritmi avansați de analiză video.

Agentic AI reprezintă însă un salt calitativ fundamental.



Dacă AI-ul tradițional funcționează ca un **instrument specializat** pe care operatorul îl activează la nevoie, **Agentic AI** operează ca un **agent autonom** care identifică singur necesitatea acțiunii și o execută.

Diferența este similară cu cea dintre un sistem care îți semnalează o anomalie pe monitor (**AI clasic**) și un sistem care detectează anomalia, evaluează severitatea, activează protocoalele adecvate și te informează doar pentru confirmare (**Agentic AI**).

În securitatea fizică, implementarea **Agentic AI** înseamnă sisteme cu capacitate de percepție contextuală, raționament strategic și acțiune autonomă pentru mitigarea riscurilor - cu intervenție umană minimă, concentrată pe supervizare și validare.

Iată cum interpretăm cele 3 capabilități fundamentale ale Agentic AI din perspectivă operațională, plus o a 4-a capabilitate critică pe care o consider obligatorie în domeniul nostru de activitate.

Agentic AI în Securitatea Fizică: De la „Instrument” la „Partener”

Agentic AI reprezintă trecerea de la un sistem care doar semnalează anomalii la un „agent” care identifică, evaluează și execută protocoale de intervenție în mod autonom. Această tehnologie integrează percepția contextuală și raționamentul strategic pentru a minimiza intervenția umană.



CELE 4 CAPABILITĂȚI FUNDAMENTALE



MANAGEMENTUL RISCULUI ȘI PRUDENȚĂ



© NotebookLM

1

Autonomous Business Operations (Operațiuni de Securitate Autonome)

În modelul tradițional de securitate fizică, procesul operațional urmează lanțul:

Detectie -> **Notificare** -> **Intervenție Umană** -> **Evaluare** -> **Răspuns**

Agentic AI scurtcircuitează acest lanț. El are capacitatea de a rula procese „capcoadă”, deoarece nu doar că detectează problema, ci inițiază și protocolul de răspuns.



Aplicabilitate în Securitate:

Vorbim despre automatizarea răspunsului la incidente în cadrul unui **SOC** (Security Operations Center) sau prin sisteme **PSIM** (Physical Security Information Management). Un agent AI poate corela date din control acces, video și efracție pentru a executa Proceduri Operaționale Standard (**POS** - Standard Operating Procedures).

Exemplu (Logistică & Depozite):

Într-un centru logistic, un senzor termic detectează o creștere de temperatură într-o zonă de depozitare a bateriilor (risc de incendiu). **Agentic AI** trimite o alarmă operatorului (care poate fi la pauză sau atent la alt monitor) și, în același timp:

1. Oprește automat sistemul de ventilație în acea zonă pentru a nu alimenta focul.
2. Deblochează turnicheții pentru evacuare.
3. Trimite coordonatele exacte către echipa de intervenție pe stațiile mobile.
4. Toate acestea se întâmplă în prima secundă, înainte ca operatorul uman să valideze alarma.



Analiza de Risc/Nota Prudentă:

Riscul de "False Positive Execution"

Dacă sistemul interpretează greșit o situație, acțiunea autonomă poate crea pagube.

Exemplu: Blocarea automată a ușilor (Lockdown) la o alarmă falsă de "active shooter" poate crea panică sau bloca accesul serviciilor de urgență.

Sfatul meu: Autonomia trebuie gradată. Începem cu **"Human-in-the-loop"** (omul aprobă acțiunea) înainte de a trece la **"Human-on-the-loop"** (omul poate valida acțiunea).

Aici trecem de la **reacție la planificare**. Managerii de securitate și consultanții de securitate se bazează adesea pe experiență și intuiție. **Agentic AI** aduce capacitatea de a simula scenarii complexe și de a propune strategii multi-step bazate pe date, nu pe "feeling".



Aplicabilitate în Securitate:

Utilizarea **Digital Twins** (Gemenii Digitali) ai clădirilor sau perimetrelor pentru a simula breșe de securitate sau fluxuri de evacuare. **Agentul AI** poate rula mii de scenarii de atac pentru a găsi vulnerabilitățile din proiectare.



Analiza de Risc/Nota Prudentă:

Dependența de Algoritm (Algorithmic Bias)

Calitatea recomandărilor AI depinde direct de calitatea datelor de antrenament.

Dacă setul de date istorice conține distorsiuni sistematice. **De exemplu**, dacă incidentele sunt înregistrate preponderent într-o anumită zonă nu pentru că acolo s-au concentrat efectiv, ci pentru că doar acolo au fost raportate, algoritmul va genera strategii care ignoră riscurile din zonele sub-raportate.

Orice strategie propusă de AI necesită validare obligatorie din partea unui consultant cu experiență operațională relevantă.



3

Self Improving Systems (Sisteme cu Auto-Îmbunătățire)

Aceasta este promisiunea supremă a Machine Learning-ului: un sistem care devine mai bun pe măsură ce este folosit.

Spre deosebire de software-ul tradițional, ale cărui performanțe sunt fixate de la momentul implementării, un sistem **ML Agentic** evoluează organic prin expunere la date reale și feedback operațional. Fiecare decizie, fiecare alarmă, fiecare interacțiune cu operatorul devine o lecție care rafinează algoritmul.

Această capacitate de adaptare continuă transformă sistemul dintr-un instrument static într-un partener dinamic care se maturizează odată cu mediul pe care îl protejează. **În securitate, stagnarea înseamnă vulnerabilitate, deoarece amenințările evoluează, iar sistemele noastre trebuie să țină pasul.** Atacatorii inovează constant, își adaptează tacticile, explorează noi vectori de intruziune, profită de schimbările din mediul fizic. Un sistem de securitate calibrat perfect astăzi va fi suboptimal peste șase luni și potențial ineficient peste un an.



Aplicabilitate în Securitate:

Reducerea alarmelor false ("pacostea") industriei de securitate, și nu numai a securității fizice.

Un sistem clasic de detecție perimetrală va alarma de fiecare dată când o vulpe trece gardul, până când îi reduci sensibilitatea, riscând să nu prinzi hoțul real. Un sistem Agentic învață feedback-ul operatorului.

Exemplu (Infrastructură Critică/Rezidențial):

Un sistem de supraveghere perimetrală generează o alarmă.

Operatorul o marchează ca "Falsă - Animal".

Sistemul nu numai că înregistrează, ci își ajustează parametrii rețelei neurale pentru a recunoaște semnătura specifică a acelui animal în condițiile respective de iluminare.

După 50 de astfel de iterații, sistemul nu mai deranjează operatorul pentru vulpi, dar rămâne alert pentru oameni.

Se "auto-tunează" constant.





Analiza de Risc/Nota Prudentă:

Data Poisoning (Otrăvirea Datelor)

Un atacator sofisticat poate exploata această funcție.

Poate genera intenționat alarme false, "sigure", într-un anumit punct pentru a "învăța" sistemul să ignore acea zonă sau acel tip de mișcare, creând o porțiță invizibilă pentru un atac real ulterior.

4

Agentic Compliance & Audit (Conformitate și Audit Agentic)

Am adăugat această capacitate pentru că, am considerat că în securitate, nu contează doar eficiența, ci și legalitatea (GDPR, drepturile omului, procedurile interne).

Un sistem de securitate, oricât de performant în detectarea amenințărilor, devine o sursă de risc juridic și reputațional dacă operează sistematic în detrimentul drepturilor persoanelor sau cu nerespectarea normelor de protecție a datelor. În contextul european actual, operatorii de sisteme de securitate se confruntă cu un cadru legislativ complex și exigent: **GDPR** impune termene stricte de retenție a înregistrărilor, iar prelucrarea datelor biometrice este condiționată de justificări legale clare și documentate.

Procedurile interne, de la gestionarea autorizărilor de acces până la protocoalele de export al datelor, nu reprezintă simple formalități contractuale sau cerințe de conformitate ISO, ci constituie mecanisme esențiale de protecție împotriva utilizării neautorizate sau disproporționate a tehnologiei.

Un sistem Agentic AI care monitorizează conformitatea în timp real nu substituie responsabilitatea umană, ci adaugă un nivel suplimentar de verificare continuă, prin care obligațiile legale abstracte sunt traduse în controale operaționale concrete. Rezultatul este dublu: pe de o parte, organizația își reduce semnificativ expunerea la sancțiuni GDPR, amenzi care pot atinge 4% din cifra de afaceri globală, iar pe de altă parte, sunt protejate drepturile fundamentale ale persoanelor aflate sub supraveghere.



Aplicabilitate în Securitate:

Un agent AI care "supraveghează supraveghetorii" și sistemele, verifică dacă înregistrările video sunt șterse conform termenelor legale, dacă accesul la datele biometrice este autorizat și dacă agenții de pază respectă ruta de patrulare conform SLA (Service Level Agreement).



Exemplu (Retail & Office):

Sistemul detectează că un operator **VSS** (Video Surveillance Systems) a făcut zoom nejustificat pe o persoană, sau a exportat date pe un stick USB neautorizat.

Agentul de conformitate blochează acțiunea și notifică automat Managerul de Securitate și DPO-ul (Data Protection Officer).

Analiza de Risc/Nota Prudentă:

Eroarea de context

Un agent de conformitate ar putea bloca o acțiune legitimă într-o situație de criză (ex: un operator exportă rapid o imagine pentru poliție în timpul unui jaf, dar agentul blochează transferul pentru că nu s-a completat formularul digital).

Flexibilitatea în criză rămâne un atribut uman.



Agentic AI în Securitate: De la Observator la Partener Activ

Schimbare de Paradigmă



Tranziția de la martor pasiv la partener activ

Agentic AI transformă securitatea dintr-un simplu observator într-un participant care acționează direct.



Provocarea cedării controlului

Implementarea necesită ca liderii să fie pregătiți să delegheze autoritatea către agenți autonomi.

Fundamentele Implementării Sigure



Rescrierea procedurilor și a răspunderii

Tehnologia nouă impune o actualizare completă a protocoloalelor de securitate și a regulilor de responsabilitate.



Omul rămâne „on-the-loop”

Până la atingerea încrederii de 100%, supravegherea umană este obligatorie pentru siguranța sistemului.

AI AGENTS

Executanții Digitali și Automatizarea Fluxurilor

Să clarificăm o confuzie comună: diferența dintre un **script simplu** și un **AI Agent**. Un **script** execută *exact* ce i-ai spus ($A + B = C$).

Un **AI Agent** are capacitatea de a interpreta cerința și de a folosi diverse unelte digitale pentru a ajunge la rezultat, adaptându-se la mici variații.



În securitatea fizică, **AI Agents** sunt cei care leagă insulele de tehnologie (VSS, Control Acces, Efracție, HR, Ticketing) care, în mod tradițional, nu vorbesc între ele.

Ei sunt "lipiciul" operațional.

Următoarea secțiune explorează cele 3 capacități fundamentale ale AI Agents din perspectivă operațională, la care am adăugat o a 4-a capacitate critică pentru protecția cibernetică a ecosistemului de securitate fizică.

AI Agents: „Lipiciul” Digital care Transformă Securitatea Fizică

Unifică sistemele fragmentate și optimizează operațiunile de securitate.

Schimbarea de Paradigmă: De la Script la Agent

SCRIPT (Executant)

Regula:
 $A + B = C$
(Execuție Rigidă)

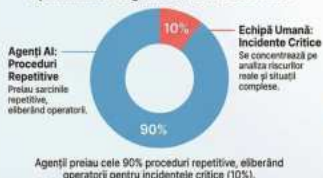
Incapabili să se adapteze la variații; necesită intervenție umană constantă.

AGENT AI (Interpret)



Interpretează cerințe, folosește unelte digitale pentru a se adapta la variații.

Optimizarea Reguli 90/10 în Securitate



Analiza Comparativă a Riscurilor Asociate Automatizării Avansate

Capabilitate	Exemplu de Risc	Măsură de Prudență
Automatizare	Conținut „Zombie” (serviciu API)	Necesitatea auditului manual periodic
Execuție API	Efectul de Cascadă (demonstrare falsă)	Validarea riguroasă a parametrilor
Monitorizare	Mediu de lucru toxic	Sisteme de toleranță pentru erori umane

Capabilități Cheie și Execuție

„Lipiciul” Sistemelor Fragmentate



Automatizarea Fluxurilor Cross-Departament (Offboarding)

Exemplu: Dezactivarea instantanee a accesului RFID și auto-plecare a unui angher.

Orchestrarea Răspunsului în Timp Real



Agentul de Igienă Cibernetică (IoT)



1

Process Automation (Automatizarea Proceselor Repetitive)

Studiile operaționale demonstrează că activitatea în securitatea fizică se împarte asimetric: aproximativ 90% constă în proceduri standardizate, repetitive, iar doar 10% în răspuns la incidente reale sau situații neașteptate.

Această majoritate procedurală, compilarea rapoartelor zilnice, procesarea cerințelor de acces, actualizarea evidențelor de personal, verificarea conformității procedurilor, reprezintă un consum masiv de ore-om pentru taskuri cu valoare adăugată redusă.

AI Agents automatizează acest segment: extrag și structurează automat datele pentru rapoarte de tură, verifică eligibilitatea vizitatorilor prin consultare multi-sistem (HR, liste de interdicție, istoric acces), mențin sincronizarea bazelor de date între VSS, control acces și ticketing, și escaladează către operator uman doar discrepanțele sau situațiile care ies din pattern-urile predefinite.

Impactul este fundamental: operatorii eliberați de povara sarcinilor administrative pot deveni specialiști concentrați pe ceea ce contează cu adevărat, analiza și gestionarea riscurilor.



Aplicabilitate în Securitate:

Automatizarea fluxurilor de date dintre departamente. De ce să introducem manual datele unui nou angajat în sistemul de control acces (PACS), când un agent AI poate "cita" baza de date HR și poate emite automat drepturile de acces bazate pe rolul angajatului?

Exemplu (Corporate Security & HR):

Procesul de **Offboarding** (plecarea unui angajat).

În momentul în care HR-ul marchează un angajat ca **"Concediat"** în sistemul lor, **Agentul AI** detectează schimbarea și execută instantaneu:

1. Dezactivarea cardului de acces fizic.
2. Anularea permisiunilor de parcare.
3. Notificarea recepției pentru a nu permite accesul ca vizitator fără aprobare specială.

Totul în timp real, eliminând fereastra de risc dintre momentul demiterii și momentul în care managerul de securitate citește email-ul de la HR.





Analiza de Risc/Nota Prudentă:

Riscul de "Zombie Accounts"

Dacă Agentul AI pierde conexiunea cu baza de date HR sau apare o eroare de sincronizare (API error), drepturile de acces pot rămâne active pentru persoane neautorizate.

Asta înseamnă că automatizarea **nu elimină** necesitatea auditului periodic manual.

2

Tool and System Execution

(Execuție și Integrare Sisteme)

Aici vorbim despre capacitatea agenților de a "apăsa butoane" digitale.

În practică, aceasta înseamnă executarea automată de comenzi și acțiuni care, în mod tradițional, ar necesita intervenția umană prin interfețe grafice separate, un operator care comută manual între mai multe ecrane, care identifică informația relevantă, decide acțiunea, apoi deschide aplicația corespunzătoare și execută comanda.

Un **AI Agent** poate efectua această secvență în milisecunde, declanșând simultan acțiuni în multiple sisteme fără latența cognitivă și operațională umană. **"Apăsarea butonului digital" poate însemna:** deblocarea unei uși, trimiterea unui email, pornirea unei înregistrări video la rezoluție maximă, ajustarea sensibilității unui senzor, escaladarea unei alarme către autorități, sau chiar comanda unei drone de patrulare autonomă.

În securitate, avem zeci de sisteme disparate. Realitatea operațională a majorității organizațiilor este un mozaic tehnologic eterogen: VMS (Video Management System) de la un producător, ACS (Access Control System) de la altul, sistemul de detecție incendiu integrat de o terță companie, senzori perimetrali care comunică prin protocoale proprietare, etc.



Aplicabilitate în Securitate:

Orchestrarea răspunsului fizic. Un AI Agent poate lega sistemul de detectare a incendiului (BMS) de sistemul de control acces și de sistemul de adresare publică (PA).



Exemplu (Retail & Logistică):

Sistemul video analitic detectează o aglomerație suspectă la o casă de marcat (posibil incident violent sau fraudă).

Agentul AI nu doar alertează, dar și execută comenzi în alte sisteme:

1. Trimite un mesaj vocal

automat în zona respectivă: "Vă rugăm păstrați distanța".

2. Trimite o notificare pe telefonul sau pe ceasul inteligent (smartwatch) al celui mai apropiat agent de pază, nu la dispecerat.

1. Marchează timestamp-ul în sistemul de management video (VMS) pentru o regăsire rapidă ulterioară.

Analiza de Risc/Nota Prudentă:

Cascada de Erori (Ripple Effect)

Un agent configurat incorect poate declanșa o reacție în lanț.

Dacă un senzor defect raportează eronat "**Incendiu**", agentul poate declanșa, prin API, deschiderea tuturor ușilor (Fail-Safe), compromițând securitatea perimetrală a unei zone sensibile (ex: trezorerie sau server room).



3

Operational Monitoring (Monitorizarea Operațională)

Managerii de securitate nu pot fi peste tot.

Limitările biologice sunt imuabile: un manager poate supraveghea efectiv 5-7 persoane direct, procesează doar câteva fluxuri de informații simultan, și are nevoie de somn și pauze. Într-o organizație cu 50 de agenți pe 3 schimburi, 100 de camere video active și sute de senzori, volumul de date operaționale orare, depășește cu ordine de mărime capacitatea cognitivă umană. **Rezultatul:** managerii reacționează post-factum, bazându-se pe rapoarte retrospective. **Problemele mici:** un senzor cu date eronate, un agent care își scurtează ruta, o cameră cu imagine degradată, rămân nevăzute până devin incidente majore.

AI Agents pot monitoriza performanța sistemului și a oamenilor 24/7, raportând doar abaterile (Management by Exception). Un agent AI nu obosește și nu are puncte oarbe cognitive; el urmărește simultan toate variabilele operaționale, poziții GPS, parametri tehnici ai camerelor, rate de alarmă, timpii de răspuns, comparându-le continuu cu baseline-ul sau pragurile SLA. Conceptul "**Management by Exception**" este esențial: în loc să bombardeze managementul cu rapoarte exhaustive despre tot ce funcționează normal, agentul filtrează zgomotul și escaladează DOAR anomaliile semnificative.



Aplicabilitate în Securitate:

Monitorizarea "sănătății" echipamentelor (Health Monitoring) și a performanței agenților de pază (Guard Tour performance).

AI Agents verifică continuu starea tehnică a infrastructurii: camerele video pentru pierderi de semnal, obturări sau degradare imagine, senzori perimetrali pentru anomalii în transmisia datelor, sisteme de control acces pentru erori de comunicare sau blocaje mecanice.

Simultan, monitorizează performanța patrulelor umane prin GPS/NFC: respectarea rutelor predefinite, timpii de staționare la punctele de control, devierile de traseu, și aderența la programul de patrulare conform SLA-urilor stabilite contractual.

Exemplu (Pază Umană și Dispecerat):

Un agent AI monitorizează în timp real rutele de patrulare prin GPS/NFC.

Dacă un agent de securitate întârzie mai mult de 5 minute la un punct de control sau deviază de la traseu, sistemul nu așteaptă raportul de dimineață. El trimite o alertă "Soft" către dispecer ("Verifică agentul X") și, dacă nu primește răspuns, escalează către Supervisor. De asemenea, monitorizează dacă, camerele video, au imaginea obturată sau pierde semnalul ("Video Loss").





Analiza de Risc/Nota Prudentă:

Micromanagement și Uzura Morală

Utilizarea AI pentru monitorizarea strictă a oamenilor poate duce la un mediu de lucru toxic și la demisia.

AI-ul trebuie setat să tolereze mici abateri umane naturale, altfel devine un instrument de teroare, nu de securitate.

4

Cyber-Physical Hygiene Agent (Agentul de Igienă Cibernetică)

Sistemele moderne de securitate fizică au suferit o metamorfoză tehnologică fundamentală, evoluând dintr-o colecție de dispozitive analogice izolate într-un ecosistem digital interconectat.

Fiecare componentă, camerele de supraveghere cu rezoluție 4K, controlerile biometrice de acces, detectoarele de mișcare cu analiză termică sau interfonul video, reprezintă de fapt un computer specializat, echipat cu unitate de procesare, memorie RAM, sistem de operare embedded (Linux, Android, sau variante proprietare) și interfață de rețea Ethernet sau wireless.

Această **convergență între securitatea fizică și infrastructura IT** amplifică considerabil eficiența operațională: management centralizat, actualizări firmware remote, integrare cu platforme cloud, analiză video bazată pe AI la sursă, dar, în același timp, suprafața de atac cibernetic se extinde exponențial.

Realitatea pieței arată dispozitive comercializate cu parole implicite neschimbate ("admin/admin"), protocoale de comunicare necriptate, firmware-uri neactualizate care conțin CVE-uri (Common Vulnerabilities and Exposures) publice de ani de zile, și absența oricăror mecanisme native de detectare a intruziunilor sau comportamentului anomal.



Aplicabilitate în Securitate:

Un agent specializat care scanează constant rețeaua de securitate (VLAN-ul de securitate) pentru a identifica vulnerabilități: firmware neactualizată, parole default neschimbate, porturi deschise inutil.



Exemplu (Smart Building/Office):

Agentul AI detectează că o cameră IP instalată recent are încă parola de fabrică ("admin/1234").

Imediat, izolează camera în rețea (o pune în carantină) și deschide un tichet automat către echipa tehnică pentru remediere, prevenind transformarea camerei într-un punct de acces pentru hackeri (Botnet).

Analiza de Risc/Nota Prudentă:

Falsa Senzație de Siguranță

Faptul că un agent AI face update-uri de firmware nu înseamnă că rețeaua este impenetrabilă.

Securitatea cibernetică necesită o abordare **"Defense in Depth"**, iar agentul este doar un strat, nu soluția totală.



Agentic AI: De la Strategie la Execuție Automatizată



STRATEGIE
AGENTIC AI

Trecerea de la Strategie la Execuție.

Agentic AI concepe planul, în timp ce AI Agents asigură coordonarea operațională rapidă.



IMPACT MULTIPLICAT AL ERORILOR

Instrucțiunile ambigue sau erorile de configurare se propagă instantaneu în tot sistemul.



Unificarea ecosistemelor fragmentate.

Transformă platformele izolate în ecosisteme collaborative unificate prin integrare cross-system.

EXECUȚIE OPERAȚIONALĂ
AUTOMATIZATĂ



VALIDAREA RIGUROASĂ ESTE IMPERATIVĂ

Testarea în medii controlate este obligatorie înainte de orice implementare în producție.



PRODUCȚIE

GENERATIVE AI

Arhitectul Culturii de Securitate și “Oracolul” Procedural

În securitatea fizică, paradoxul este cunoscut: **investim în tehnologie de vârf, dar majoritatea breșelor provin din eroare umană.**

Angajatul care permite accesul fără verificare, operatorul care ignoră o alertă considerând-o "probabil falsă", utilizatorul care ocolește procedura pentru "a câștiga timp", acestea sunt vulnerabilitățile reale.

Problema de fond nu este rea-voință, ci două deficiențe structurale: **(1) conștientizarea superficială** - training-uri generice, odată pe an, care nu generează schimbare comportamentală reală, și **(2) documentație impracticabilă**, proceduri de securitate voluminoase, formulate birocratic, pe care nimeni nu le citește efectiv în momentul critic.



Generative AI (GenAI) devine soluția acestei rupturi. Nu este vorba doar despre chatbot-uri sau generatoare de text, ci despre capacitatea de a crea **conținut adaptat dinamic**: explicații personalizate pentru proceduri complexe, scenarii de training contextualizate pe rolul specific, ghiduri vizuale generate instant pentru situații neobișnuite. În loc de manuale statice de 200 de pagini, utilizatorul primește răspunsuri precise, în limba sa, la întrebarea sa specifică, în momentul în care are nevoie.

De altfel, acest white paper ,în sine este un exemplu concret de utilizare a Generative AI, transformând concepte tehnice complexe din surse internaționale în conținut structurat, accesibil și adaptat specificului industriei de securitate fizică din România.

GenAI: Revoluționarea Culturii și Procedurilor de Securitate Fizică

Deși companiile investesc în tehnologie de vârf, majoritatea breșelor de securitate provin din eroare umană și proceduri imposibile de parcurs în momente critice. GenAI acționează ca un "Oracol" procedural, transformând manualele aride în dialoguri interactive și instruirea generică în conștientizare personalizată.

Problema: Barierele Securității Tradiționale



Paradoxul Tehnologiei vs. Eroarea Umană
Majoritatea breșelor apar din ignorarea alertelor sau ocolirea procedurilor pentru a câștiga timp.



Documentație Impracticabilă
Manualele statice de sute de pagini sunt imposibile de consultat eficient în situații de criză.



Conștientizare Superficială
Training urle anuale generice nu reușesc să schimbe comportamentul real al angajaților.



De la Reguli Aride la "Awareness" prin Poveste
Transformă restricțiile în campanii personalizate pe rolul angajatului (recepție, IT, management).



"Oracolul" Procedural (RAG)
Oferă răspunsuri instantanee din biblioteca de proceduri în timpul incidentelor critice.



Generarea de Date Sintetice
Creează scenarii vizuale fotorealiste pentru antrenarea algoritmilor de detecție fără riscuri reale.

COMPARAȚIE IMPACT OPERAȚIONAL

Abordare Tradițională

2-4 săptămâni / Costuri mari
Intervenție umană repetitivă
Minute prețioase (navigare foldere)
Căutare Proceduri

Impact GenAI

5-10 minute / Cost zero
Creare Training
Automatizare 70-80% (Chatbot)
Secunde (dialog în limbaj natural)

Managementul Riscului și Responsabilitatea



AI-ul creează, Omul validează
GenAI poate "halucina"; expertul uman rămâne singurul responsabil pentru verificarea acurateții.



Riscul "Data Leakage"
Utilizarea instrumentelor publice poate expune planuri de cădrări sau date confidențiale către terți.

NotebookLM

1

Marketing and Content Creation

(Marketing și Creare de Conținut)

În domeniul securității fizice, echivalentul "marketingului" este **Security Awareness** (procesul de educare și conștientizare a utilizatorilor cu privire la riscuri și comportamente corecte).

Realitatea este că profesioniștii în securitate au dezvoltat o reputație nemeritată pentru **comunicare defensivă**: instrucțiuni formulate ca restricții ("Este interzis să..."), documentație aridă, ton autoritar și un limbaj tehnic și/sau de specialitate, uneori, inaccesibil persoanelor cărora li se adresează.

Rezultatul este previzibil, managementul și angajații, văd securitatea ca pe un set de obligații birocratice, nu ca pe o protecție pentru ei înșiși. **Generative AI (GenAI)** schimbă fundamental această dinamică.

În loc să **impunem** reguli, putem să **explicăm** riscurile, în loc să **interzicem** comportamente, putem să **demonstrăm** consecințele.

GenAI permite managerilor de securitate să transforme comunicarea de la "**compliance prin teamă**" la "**înțelegere prin poveste**", exact ceea ce înseamnă să "vinzi" securitatea ca valoare, nu ca restricție.



Aplicabilitate în Securitate:

Dezvoltarea rapidă de module de instruire personalizate care elimină șablonismul generic și contraproductiv.

În loc de materiale universale care nu rezonază cu nimeni specific, **GenAI** permite crearea de campanii adaptate la: **(1) rolul angajatului** (recepție vs. IT vs. management), **(2) cultura organizațională** (startup tech vs. instituție bancară), **(3) incidentele istorice** specifice companiei, și **(4) limbajul și contextul local**. **Rezultatul:** conștientizare care generează schimbare comportamentală reală, nu doar bifarea unei cerințe de conformitate.

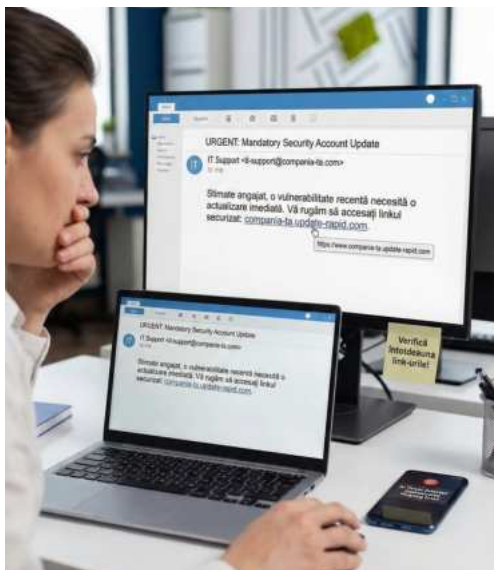
Exemplu (Corporate & Training):

Organizația dorește să lanseze o campanie de conștientizare despre riscurile de Social Engineering la nivel de recepție și acces vizitatori.

În **abordarea tradițională**, acest demers ar implica: contractarea unei agenții externe, brief-uri multiple, iterații de revizuire, producție video profesională, un proces de 2-4 săptămâni și un buget de câteva mii de euro.

Cu **Generative AI**, managerul de securitate formulează un **prompt structurat**: "Creează un scenariu de joc de rol de 2 minute între un recepționar și un presupus curier care încearcă să obțină acces neautorizat folosind tehnici de manipulare psihologică. Include un poster vizual A3 cu cele 5 red flags ale Social Engineering specifice recepției, plus 3 email-uri de teasing pentru distribuire cu 3 zile înainte de training, care să genereze curiozitate fără a dezvălui tema."





Analiza de Risc/Nota Prudentă:

Phishing-ul Perfecționat

Aceeași unealtă pe care o folosim noi pentru apărare este folosită și de atacatori.

GenAI poate crea email-uri de phishing perfecte din punct de vedere gramatical și contextual, făcând distincția dintre un email legitim și unul malițios extrem de dificilă pentru angajați

2

Customer Support Automation (Automatizarea Suportului pentru Clienți/Angajați)

În arhitectura funcțională a majorității organizațiilor, departamentul de securitate funcționează, în fapt, ca un centru de asistență intern de specialitate.

Similar unui **birou de asistență IT** care rezolvă probleme tehnice, echipa de securitate gestionează un volum constant de solicitări procedurale, întrebări recurente și cereri de clarificare.

Analiza **solicităților înregistrate** într-un departament mediu de securitate corporativă relevă un **tipar consistent**: aproximativ 70-80% din interacțiunile cu angajații sunt cereri standard, repetitive, pierderi de **cartele de acces**, clarificări despre politici de parcare, întrebări despre procedura de autorizare vizitatori, solicitări pentru accesul temporar în zone restricționate, întrebări despre protocoale de urgență.

Această realitate creează o **tensiune operațională**: resursele umane dedicate, operatori de dispecerat, agenți de securitate, manageri, petrec timp semnificativ răspunzând la întrebări cu răspunsuri standardizate, în detrimentul sarcinilor care necesită cu adevărat judecată umană și experiență de specialitate.



Aplicabilitate în Securitate:

Chatbot-uri interne de securitate (**Security Helpdesk**) antrenate pe politicile, procedurile și specificul organizațional, care oferă răspunsuri instantanee la întrebările procedurale recurente.

Impact triplu: (1) angajații obțin informații imediate, fără timp de așteptare, (2) răspunsurile sunt uniforme și actualizate conform politicilor în vigoare, (3) operatorii umani sunt disponibili pentru sarcini critice, monitorizare activă, analiză de anomalii, răspuns la incidente. **Conversația este naturală**, contextul este păstrat între întrebări, și sistemul poate solicita clarificări ("La care dintre cele 3 clădiri vă referiți?") sau transfera automat către un operator uman când detectează cerințe complexe.



Exemplu (Travel Security):

Un angajat urmează să plece într-o delegație într-o zonă cu risc ridicat.

În loc să sune managerul de securitate, interoghează bot-ul intern: *"Care sunt riscurile de securitate în Mexico City și ce procedură de transport trebuie să folosesc?"*.

Bot-ul, antrenat pe politicile companiei și rapoarte de țară actualizate, îi generează instant o informare de călătorie personalizată, cu numere de urgență și lista hotelurilor aprobate.

Analiza de Risc/Nota Prudentă:

Scurgerea de Date (Data Leakage)

Riscul critic al automatizării prin **GenAI** este exfiltrarea neintenționată de informații sensibile către sisteme externe.

Scenariul real: angajați bine-intenționați care, pentru "eficiență", utilizează instrumente **GenAI publice și gratuite** (*ChatGPT, Claude, Gemini în variantele lor publice*) pentru a formula răspunsuri la situații de securitate, introducând involuntar date confidențiale:

"Am un vizitator de la [nume client major] care vine mâine. Cum completez formularul de acces pentru el?". "Poți să-mi generezi un raport despre incidentul din [locație specifică] de aseară?". "Rezumă-mi politica noastră de acces pentru zona [infrastructură critică]" .



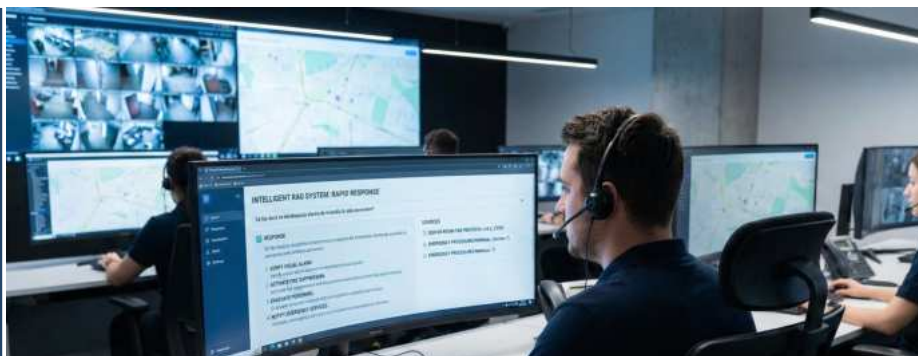
3

Internal Knowledge Assistants (Asistenți Interni de Cunoștințe –"Oracolul")

Aceasta este **provocarea fundamentală a oricărui sistem complex de securitate**: documentația exhaustivă care devine inaccesibilă exact când este vitală. Organizațiile mature au dezvoltat de-a lungul anilor bibliorafturi întregi de proceduri operaționale standard (POS), planuri de răspuns la incidente, protocoale de urgență, instrucțiuni tehnice, documente meticuloase elaborate, revizuite, aprobate.

Problema nu este lipsa documentației, ci imposibilitatea găsirii informației corecte în timp util.

Realitatea operațională: în momentul critic (alarmă de efracție la 3 dimineața, incendiu la subsol, pană de curent în data center) operatorul trebuie să navigheze printr-un labirint de foldere, să identifice manualul corect (versiunea actualizată, nu cea din 2019), să găsească capitolul relevant, să extragă pașii specifici situației. Acest proces consumă minute prețioase și introduce riscul imens al erorii umane sub stres, ceea ce înseamnă omiterea unui pas critic sau aplicarea procedurii greșite pentru situația dată. Investim zeci de ore în scrierea procedurilor perfecte, dar când sunt cu adevărat necesare, nu pot fi accesate eficient.



Aplicabilitate în Securitate:

Sisteme inteligente de recuperare și generare de răspunsuri (tehnic: RAG - Retrieval-Augmented Generation, adică Recuperare Augmentată prin Generare).

Esența tehnologiei: în loc să cauți manual prin documente, "dialoghezi" direct cu întreaga ta bibliotecă de cunoștințe. **Funcționarea simplificată:** sistemul indexează automat toate procedurile, manualele, protocoalele organizației. Când operatorul pune o întrebare în limbaj natural ("Ce fac dacă se declanșează alarma de incendiu în sala serverelor?"), sistemul:

1. **Identifică** documentele relevante din întreaga bibliotecă
2. **Extrage** pasajele specifice din acele documente
3. **Generează** un răspuns coerent, structurat, cu pași exacți
4. **Include** sursa (link către documentul original) pentru verificare

Beneficiul pentru operatorii din centrele de supraveghere și comandă: informația critică în secunde, nu minute. Fără navigare prin foldere, fără căutare în index, fără riscul de a consulta versiunea învechită a procedurii.

Exemplu (Crisis Management & SOC):

Ora 03:00 dimineața. Alarmă de inundație la subsolul 2.

Operatorul nou angajat se panichează. În loc să răsfoiască manualul, tastează sau întreabă vocal: "Procedură inundație subsol servere".

Asistentul **GenAI** extrage instantaneu pași exacți:

1. **Nu intrați (risc electrocutare).**
2. **Sunați la Mentenanță (Nr: 07xx...).**
3. **Notificați IT Manager (Nr: 07yy...).**

Îi livrează informația *critică*, filtrând zgomotul.





Analiza de Risc/Nota Prudentă:

Halucinațiile AI

GenAI poate inventa fapte cu o încredere debordantă.

Dacă procedura nu există sau e ambiguă, AI-ul ar putea inventa un pas periculos (ex: *"Opreți pompa principală"* – când de fapt, trebuia pornită).

Verificarea umană a sursei
([link către pagina din manual](#))
este obligatorie.

4

Syntetic Data Generation (Generarea de Date Sintetice)

Algoritmii inteligenți de detecție video funcționează prin învățare din exemple.

Pentru ca o cameră video să recunoască automat o armă, sau un comportament violent, trebuie "antrenată" pe mii de imagini și secvențe video care prezintă exact aceste situații, în diverse condiții de lumină, unghiuri, poziții, tipuri de amenințări.

Paradoxul este evident: de unde luăm aceste mii de exemple? Din fericire, organizațiile responsabile nu au arhive extinse cu jafuri armate, agresiuni sau intruziuni reale în propriile locații, **lipsa acestor evenimente este tocmai scopul securității preventive**, iar materialele publice disponibile (*instruire forțe de ordine, cazuri judiciare*) sunt insuficiente numeric și nu reflectă specificul fiecărei locații: geometria spațiului, tipul de iluminare, contextul operațional particular.

Rezultatul: sistemele inteligente de securitate nu au acces la volumul necesar de date reale pentru a fi antrenate eficient. Lipsa datelor devine bariera tehnologică, chiar dacă această lipsă este, ironic, un semn al succesului măsurilor de securitate existente.



Aplicabilitate în Securitate:

Folosirea **GenAI** pentru a crea imagini și clipuri video fotorealiste cu situații de risc (oameni cu arme, incendii, intrări prin efracție) pentru a antrena și testa sistemele de analiză video, fără a fi nevoie să simulăm fizic aceste scenarii.

Exemplu (R&D & System Testing):

Înainte de a instala un sistem scump de detecție a armelor, **consultantul de securitate** poate cere generarea a 500 de imagini cu diverse tipuri de arme ascunse parțial, în condiții de iluminare specifice locației clientului, pentru a testa rata de detecție a sistemului propus (**Proof of Concept digital**).





Analiza de Risc/Nota Prudentă:

Overfitting (Suprareglare) pe Date Sintetice

Dacă antrenăm sistemul doar pe date generate de AI, s-ar putea să nu recunoască realitatea "murdară" și imperfectă.

Datele sintetice trebuie să completeze, nu să înlocuiască datele reale.

AI Generativ în Securitate: Amplificator, nu Înlocuitor

AMPLIFICAREA EXPERTIZEI UMANE

Amplificator dual: Voce și Memorie



GenAI extinde capacitatea de comunicare personalizată și oferă acces instantaneu la întreaga cunoaștere organizațională.



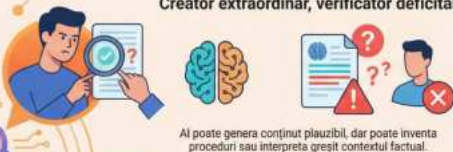
Impact multiplicat, nu înlocuit

Rolul AI este de a asista specialistul să comunice eficient cu sute de angajați simultan.



LIMITELE ȘI RESPONSABILITATEA EXPERTULUI

Creator extraordinar, verificador deficitar



Fără transfer de responsabilitate

Responsabilitatea pentru validare și consecințe aparține integral expertului uman, nu instrumentului tehnologic.



"Validarea este esența profesiei"

Folosii AI pentru educație și comunicare, fără a abdică de la expertiza care vă definește.

NEURAL NETWORKS

“Simțurile” Sistemului și Recunoașterea Modelelor

Rețelele neuronale reprezintă arhitectura fundamentală care face posibilă viziunea computerizată (**Computer Vision**).

Evoluția tehnologică este profundă: în modelul tradițional, o cameră de supraveghere funcționa ca un simplu captator optic; înregistra și transmitea imagini către operatorul uman care avea sarcina de interpretare. Camera în sine era "oarbă", adică detecta schimbările de pixeli (mișcarea), dar nu înțelegea ce reprezintă acei pixeli.



Rețelele neuronale moderne transformă fundamental această paradigmă: camera nu mai este doar un senzor pasiv, ci devine un sistem de percepție inteligent.

În loc să identifice doar că "s-a mișcat ceva în cadrul 5", sistemul recunoaște și clasifică: "persoană umană", "vehicul tip autoturism", "obiect metalic cu formă de armă", "flacără și fum - posibil incendiu", "model de mișcare agresivă".

Camera dobândește capacitatea de înțelegere contextuală a scenei vizuale.

Această capacitate redefineste utilitatea operațională a sistemelor de supraveghere video: trecerea de la instrumente de analiză post-factum (bune pentru investigații după incident, identificare suspecti, reconstituire evenimente) la instrumente de prevenție activă (capabile să detecteze amenințări în dezvoltare și să declanșeze alerte înainte ca incidentul să se consume).

Evoluția Securității: De la Camere Pasive la Percepție Inteligentă

Rețelele neuronale transformă camerele de supraveghere din simpli senzori "orbi" în sisteme capabile să înțeleagă contextul. Această evoluție permite trecerea de la investigații post-eveniment la prevenția activă, detectând amenințări prin analiză vizuală, audio și comportamentală.

DIMENSIUNILE PERCEPȚIEI INTELIGENTE



De la Pixeli la Înțelegere Contextuală

Sistemul nu detectează doar mișcare, ci clasifică obiecte, persoane și comportamente agresive în timp real.

Înțelegere Contextuală

Sistemul nu detectează doar mișcare, ci clasifică obiecte, persoane și comportamente agresive în timp real.



Analiza Audio Omnidirecțională

Microfonul eșină auzul din toate direcțiile, detectând escaladarea verbală și pericula înaintea de violența fizică.



Biometrie cu Acuratețe de 99%+

Rețelele moderne extrag amprente faciale digitale unice, imune la schimbări de look sau unghiuri variate.

COMPARAȚIA LOGICII DE PREDICȚIE



PROVOCĂRI ȘI FACTORI CRITICI



Riscul de "Alarm Fatigue"

Duturile de antrenament deficiente generează alerte false masive, determinând operatorii umani să ignore amenințările reale.

Atacuri Adversariale și Deepfakes



Procesare "On the Edge"



1

Visual Quality Inspection (Inspecția Vizuală a Calității)

Principiul inspecției vizuale automate prin rețele neuronale este identic în manufacturing și în securitate: detecția automată a anomaliilor vizuale, dar definiția "anomaliai" diferă fundamental în funcție de domeniu.

În **contextul industrial clasic** (linii de producție, control calitate), "anomalia" înseamnă abateri de la standardul de fabricație: o zgârietură pe vopsea, o piesă lipită incorect, o sudură defectă, o etichetă aplicată strâmb. Sistemul compară produsul real cu modelul perfect și identifică discrepanțele care îl fac neconform specificațiilor.

În **securitatea fizică**, același mecanism tehnologic operează, dar "anomalia" capătă semnificație de risc: nu căutăm imperfecțiuni estetice, ci abateri de la standardele de siguranță și protocoalele de securitate.

"Defectul" devine: o breșă în perimetru (gard deteriorat, poartă lăsată deschisă), echipament de protecție absent (lucrător fără cască în zonă obligatorie), obiecte neautorizate în zone restricționate (bagaj suspect lăsat nesupravegheat), comportamente care deviază de la modelul normal (*persoană care pătrunde prin zone neobișnuite în loc de trasee standard*). **Esența rămâne aceeași: "recunoaște ce este normal, semnalează ce este diferit", dar în securitate, "diferitul" nu este neconformitate de calitate, ci indicator potențial de amenințare.**



Aplicabilitate în Securitate:

Monitorizarea automată a echipamentului de protecție și detectarea intrușilor care încearcă să se camufleze. Rețeaua neuronală este antrenată să recunoască "normalul" (oameni cu cască și vestă) și să alerteze instantaneu la "anormal" (persoană civilă în zonă restricționată).



Exemplu (Industrial & Critical Infrastructure):

Pe un șantier sau într-o rafinărie, camerele monitorizează poarta de acces auto.

Rețeaua neuronală nu citește doar numărul de înmatriculare (LPR - tehnologie veche), ci inspectează vizual vehiculul:

- ✓ Există scurgeri vizibile sub camion?
- ✓ Șoferul poartă centură?
- ✓ Încărcătura este asigurată corect?.

Dacă detectează o anomalie (ex: *un obiect suspect sub șasiu*), blochează automat bariera.

Analiza de Risc/Nota Prudentă:

Atacuri Adversariale (Adversarial Attacks)

Rețelele neuronale pot fi păcălite de modele vizuale special create (ex: un tricou cu un tipar specific care face persoana "invizibilă" pentru algoritmul de detectare a persoanelor).

Securitatea nu trebuie să se bazeze *exclusiv* pe analiză video.



2

Voice and Call Analysis (Analiza Audio și Detecția Agresiunii)

Sistemele de supraveghere video au o limitare fizică fundamentală: dependența de linia directă de vedere.

Orice cameră video de supraveghere, indiferent de rezoluție sau unghi de vizualizare, poate monitoriza doar ceea ce este vizibil optic din punctul său de montare. Acest lucru creează inevitabil **zone oarbe** (unghiuri moarte), spații ascunse în spatele coloanelor, sub mese, în colțuri, după mobilier, unde activitatea rămâne nedetectată. Un agresor care cunoaște poziționarea camerelor poate exploata aceste zone pentru a evita detectarea.

Monitorizarea audio prin microfoane operează pe un principiu fizic diferit: **sunetul nu necesită linie directă de vedere**. Undele sonore se propagă omnidirecțional (în toate direcțiile), reflectă pe suprafețe, ocolesc obstacole și pătrund în spații pe care camerele nu le pot vedea. **Un microfonul poziționat strategic poate detecta evenimente acustice** din întregul spațiu înconjurător, indiferent de obstacole vizuale. Colțul ascuns în spatele teșghelei, holul din spatele ușii, zona din spatele recepției, toate devin "audibile" chiar dacă rămân "invizibile".

În contextul securității fizice, tehnologia de analiză a sentimentelor ("*sentiment analysis*" - capacitatea de a detecta starea emoțională din voce) capătă o aplicație preventivă critică: **identificarea escaladării agresiunii înainte de tranziția la violență fizică.**



Aplicabilitate în Securitate:

Senzori audio în spații publice, spitale sau ghișee, antrenați să detecteze semnătura acustică a țițetelor, a geamurilor sparte sau a focurilor de armă (Gunshot Detection).

Exemplu (Spitale & Retail):

La unitatea de primiri urgențe (UPU), tensiunea este mare.

Un senzor audio detectează o creștere bruscă a tonului vocii și anumite cuvinte cheie asociate amenințării.

Sistemul alertează discret paza înainte de a se arunca primul pumn.

De asemenea, în dispecerat, poate analiza apelurile radio ale agenților de pază: dacă detectează stres în vocea agentului, trimite automat backup, chiar dacă agentul nu a cerut explicit ajutor (Panic Analysis).





Analiza de Risc/Nota Prudentă:

Confidențialitatea (GDPR & Privacy)

Înregistrarea conversațiilor ambientale este o zonă legislativă minată.

Soluția tehnică este procesarea "on the edge" (local, în dispozitiv), unde sistemul nu înregistrează cuvintele, ci doar analizează frecvențele și decibelii, ștergând datele imediat după analiză.

3

Personalized Recommendations (Profilare Comportamentală și Predicție)

Platformele de recomandare de conținut precum Netflix funcționează pe un principiu fundamental de învățare automată: analizează istoricul comportamental al utilizatorului pentru a prezice preferințele viitoare.

Sistemul observă ce filme ai vizionat anterior, cât timp ai stat pe fiecare, ce ai abandonat rapid, ce ai revăzut, și identifică modele statistice: *"Dacă utilizatorul a vizionat 5 filme de acțiune și 3 thrillere în ultima lună, probabilitatea că va aprecia următorul thriller disponibil este 78%"*. Recomandarea nu este ghicire, ci predicție bazată pe date istorice.

Rețelele neuronale în securitate operează pe exact același principiu tehnic: **învățare din comportament istoric pentru predicția anomaliilor**, dar cu o diferență fundamentală de aplicație și consecințe.

Esența tehnologică este identică: ambele sisteme **compară prezentul cu modele învățate din trecut și generează recomandări bazate pe probabilități statistice.**

Diferența critică: în entertainment, o recomandare greșită înseamnă un film nepotrivit; **în securitate, o "recomandare" greșită** (fals negativ) poate însemna un incident nedetectat cu consecințe grave, iar o alertă falsă (fals pozitiv) înseamnă resurse consumate inutil și oboseală de alarmă.

Din această perspectivă, rețeaua neurală de securitate nu "decide" sau "declară", ci **recomandă atenție umană** bazată pe calcul probabilistic: *"Am observat o abatere de X% de la pattern-ul normal învățat din 10.000 ore de monitorizare. Sugerez verificare umană pentru confirmare."*



Aplicabilitate în Securitate:

Detectarea comportamentului suspect (Loitering) și a tiparelor de trafic anormale. Sistemul învață fluxul normal al oamenilor într-o clădire și semnalează deviațiile.



Exemplu (Aeroporturi & Mall-uri):

Sistemul observă o persoană care a lăsat un bagaj jos și s-a îndepărtat rapid (Abandonment detection), sau o persoană care merge "contra fluxului" pe un coridor de ieșire.

Rețeaua neuronală face o "recomandare" operatorului: *"Probabilitate 85% incident de securitate la Poarta 4. Verificare necesară"*.

Nu este o certitudine, este o predicție statistică bazată pe mii de ore de antrenament.

Analiza de Risc/Nota Prudentă:

Biasul Algoritmice (Prejudecata)

Dacă datele de antrenament conțin prejudecăți (ex: corelarea anumitor tipuri de îmbrăcăminte cu infracționalitatea), sistemul va discrimina, generând alerte false pe criterii incorecte etic.

Validarea umană este crucială pentru a evita profilarea abuzivă.



4

Biometric Identification (Identificarea Biometrică)

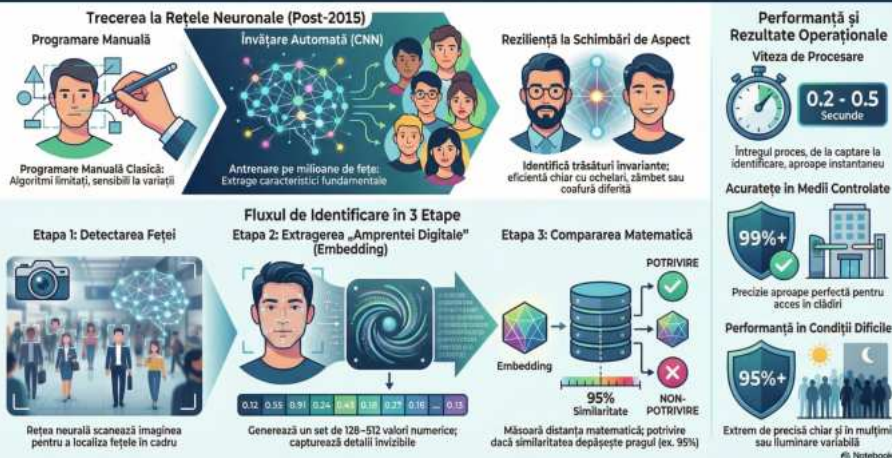
Tehnologia de recunoaștere facială (FR - Face Recognition) nu ar fi posibilă la nivel practic și scalabil fără arhitectura rețelor neuronale profunde.

Acestea reprezintă infrastructura computațională fundamentală care face ca un sistem să poată distinge între milioane de fețe diferite cu precizie ridicată și în timp real.

În trecut (**anii 2000**), sistemele de recunoaștere facială foloseau algoritmi clasici bazați pe **măsurători geometrice simple**: distanța între ochi, lungimea nasului, forma bărbiei - aproximativ 20-30 de puncte de referință (landmarks).

Aceste sisteme erau **fragile**, deoarece schimbări minore de iluminare, unghi al capului, expresia facială sau ochelarii, făceau ca sistemul să eșueze în recunoaștere. Acuratețea era de 60-70% în condiții controlate, asta fiind, practic, inutilizabil în medii reale.

Revoluția Identificării Biometrice: Cum Funcționează Recunoașterea Facială Modernă



Aplicabilitate în Securitate:

Control acces fără atingere (frictionless) și liste negre (Blacklist) automate.

Recunoașterea facială elimină fricțiunea operațională clasică a controlului acces, nu mai există carduri uitate acasă, PIN-uri compromise, sau cozi la turnicheturi în orele de vârf. **Identitatea devine inseparabilă de persoană**: autentificarea se produce pasiv, în mers, fără nicio interacțiune conștientă din partea utilizatorului. Această fluiditate nu sacrifică securitatea, ci o amplifică, spre deosebire de un card de acces, fața nu poate fi împrumutată, clonată fizic sau transferată unui terț neautorizat.

Simultan, **funcționalitatea Blacklist transformă radical paradigma de securitate reactivă**. Sistemul nu mai așteaptă ca o persoană cu intenții ostile să ajungă la recepție pentru a fi identificată manual, scanează în timp real fiecare față din câmpul vizual, comparând-o automat cu bazele de date de persoane interzise.



Exemplu (Corporate & VIP):

CEO-ul intră în clădire.

Camera de la turnichet îi recunoaște fața în 0.2 secunde, îi deschide ușa și cheamă liftul la etajul executiv.

În același timp, sistemul scanează mulțimea din lobby pentru a identifica fețe cunoscute din baza de date "Persoane Interzise" (foști angajați concediați conflictual).

Analiza de Risc/Nota Prudentă:

Deepfakes și Spoofing

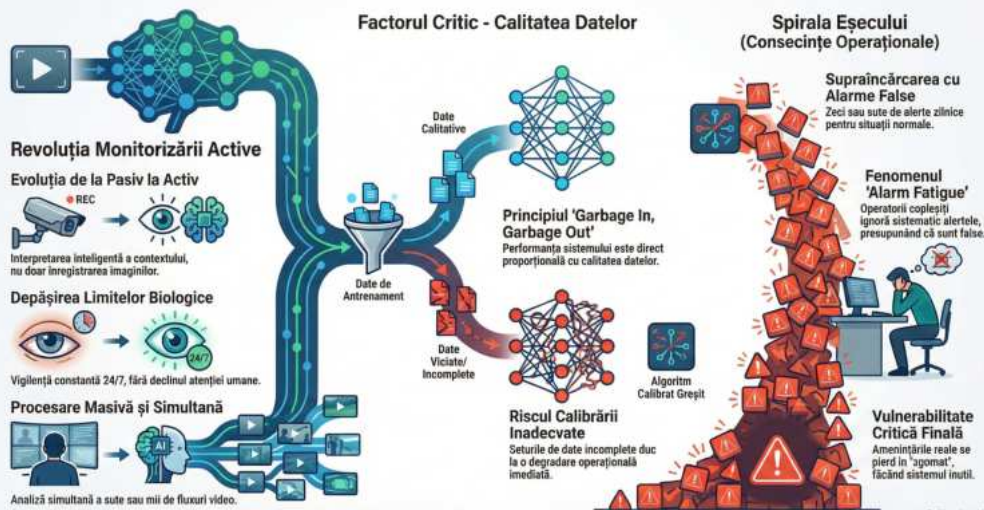
Tehnologia avansează și de partea cealaltă.

O fotografie de înaltă rezoluție sau un video "Deepfake" pe o tabletă poate păcăli o cameră simplă.

Sistemele moderne necesită "**Liveness Detection**" (verificarea faptului că este o ființă vie, 3D, nu o poză 2D) pentru a fi sigure.



Rețelele Neuronale în Securitate: Între Performanță și Riscul "Oboselii de Alarmă"



Rețelele neuronale marchează tranziția fundamentală a sistemelor de securitate de la **captare pasivă** la **percepție activă și interpretare**. Acestea **elimină limitările biologice ale operatorilor umani**: nu sunt afectate de oboseală după ore de monitorizare, mențin vigilență constantă fără declin de atenție, și pot procesa simultan sute sau chiar mii de fluxuri video, ceva imposibil pentru echipe umane indiferent de dimensiunea lor.

Însă, această capacitate tehnologică extraordinară vine cu o dependență critică: **performanța sistemului este direct proporțională cu calitatea datelor folosite pentru antrenament**. Un algoritm antrenat pe seturi incomplete, nereprezentative sau viciate va produce rezultate la fel de deficitare, principiul fundamental în inteligența artificială: "dacă introduci date proaste, vei obține rezultate proaste".

Consecința practică a unei calibrări inadecvate este degradarea operațională prin supraîncărcare cu **alarme false**: sistemul generează zeci sau sute de alerte zilnice pentru situații normale interpretate greșit ca amenințări. Operatorii, copleșiți de fals pozitive constante, dezvoltă ceea ce se numește "oboseală de alarmă" (alarm fatigue) încep să ignore alertele, presupunând că "probabil este din nou o falsă alarmă".

Rezultatul final: un sistem teoretic avansat devine practic inutil, deoarece semnalele reale de amenințare se pierd în zgomotul alarmelor false și sunt ignorate împreună cu restul.

MACHINE LEARNING

Analiza Predictivă și Optimizarea Resurselor

Securitatea fizică a funcționat istoric pe un model fundamental reactiv:

detectăm amenințarea → **răspundem la amenințare** → **analizăm post-incident.**

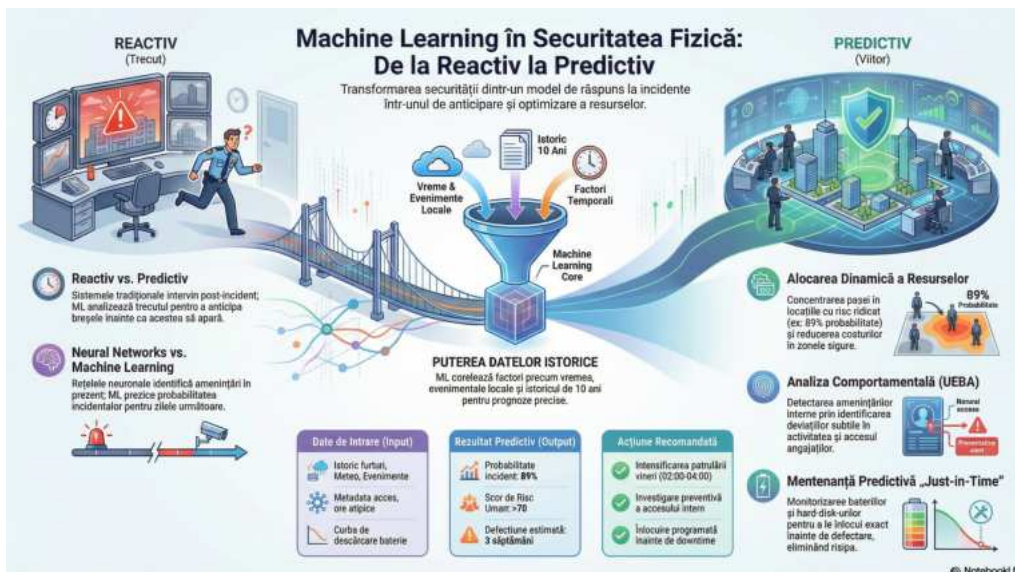
Sistemele tradiționale sunt construite să reacționeze: alarma se declanșează când ușa este forțată, camera înregistrează când intrusul este deja în clădire, paza intervine după ce breșa s-a produs. Este modelul "detectează și răspunde", eficient pentru minimizarea pagubelor, dar nu pentru prevenirea incidentului în sine.



Machine Learning (ML - Învățare Automată) schimbă fundamental această paradigmă prin introducerea dimensiunii **temporale predictive**: în loc să reacționăm la prezent, analizăm trecutul pentru a anticipa viitorul. **Diferența față de Rețelele Neuronale** (care operează în timp real) este critică.

Procesare ML: Algoritmul identifică corelații statistice: "În ultimii 10 ani, 73% din tentativele de efracție la depozite au avut loc vineri/sâmbătă noapte (2-5 AM), în perioade cu vreme rea (ploaie/ceață), când echipa de pază era redusă (programare sărbători), și când avea loc un eveniment major în oraș care distraea atenția autorităților."

leșire predictivă: "Săptămâna viitoare: vineri noapte, prognoza meteo arată ploaie, programul arată efectiv redus de pază (concedii), în oraș are loc meci de fotbal major → Probabilitate 89% tentativă efracție Depozit B, interval 02:00-04:00
→ Recomandare: intensificați patrularea, activați camere suplimentare, alertați echipa de răspuns"



1

Demand and Sales Forecasting (Predictia Cererii și a Riscului)

Tehnologia Machine Learning de predicție a cererii (demand forecasting) operează pe același principiu fundamental în business și în securitate: analiză de modele istorice pentru anticiparea nevoilor viitoare, dar cu diferențe critice în ce se prezice și ce se optimizează.

În Securitatea Fizică:

- ✓ **Ce se prezice:** Probabilitatea incidentelor (furturi, efracții, agresiuni), intensitatea necesarului de resurse de pază
- ✓ **Date analizate:** Istoric incidente, sezonabilitate (sărbători = magazine aglomerate = risc furturi crescut), evenimente locale (meciuri, festivaluri = distrageri), meteo (ploaie = vizibilitate redusă = risc efracții), condiții economice
- ✓ **Obiectivul:** Alocarea dinamică a resurselor de securitate - concentrarea pazei unde/când riscul este maxim, reducerea unde riscul este minimal
- ✓ **Output tipic:** "Săptămâna viitoare, creștere 65% risc furturi Magazin Nord → dublați efectivul de pază; risc redus Magazin Sud → reduceți la efectiv minim"



Aplicabilitate în Securitate:

Predictive Policing (Supraveghere Predictivă) aplicată în mediul privat. Alocarea dinamică a resurselor de securitate în funcție de tendințele de risc, nu după un orar fix și rigid.

Exemplu (Retail & Evenimente):

Un lanț de magazine folosește ML pentru a analiza furturile din ultimii 5 ani.

Algoritmul corelează furturile cu factori externi: sărbători legale, meciuri de fotbal în zonă, starea vremii (ploaia reduce traficul, dar crește riscul de jafuri "hit and run").

Sistemul generează o prognoză: "Săptămâna viitoare, dublați paza la magazinul din Nord, dar reduceți-o la cel din Sud".

Eficiență maximă,
costuri optimizate.





Analiza de Risc/Nota Prudentă:

Profeția Auto-Împlinită

Dacă trimiți mereu paza într-o zonă pentru că algoritmul zice că e "riscantă", vei prinde mai multe incidente acolo, ceea ce va confirma algoritmului că zona e riscantă.

Astfel, ignori alte zone care pot deveni vulnerabile pentru că nu le monitorizezi.

2

Customer Risk and Churn Prediction (Predicția Abandonului Clienților)

Tehnologia de predicție a "abandonului clienților" (churn prediction) din domeniul comercial își găsește o aplicație transformată în securitatea fizică.

În loc să identificăm clienții care vor părăsi compania, **identificăm angajați care prezintă risc de a deveni amenințări interne și personal de securitate aflat în pragul demisiei/demiterii.**

Amenințarea din interior se referă la **angajați, contractori sau parteneri cu acces legitim la clădiri, sisteme și informații care devin o amenințare intenționată sau neintenționată** pentru organizație. Exemple: furt de date confidențiale, sabotaj echipamente, facilitare acces neautorizat pentru terți, violență la locul de muncă.

Provocarea: Acești indivizi au deja acces autorizat: carduri de acces valide, parole funcționale, cunoștințe despre proceduri. Sistemele tradiționale de securitate (care blochează pe cei "din afară") sunt inefficiente împotriva amenințărilor "din interior".



Aplicabilitate în Securitate:

Analiza comportamentală a angajaților (UEBA - User and Entity Behavior Analytics). ML detectează deviațiile subtile care indică un angajat nemulțumit, gata să fure date sau să saboteze, sau un agent de pază pe cale să demisioneze (burnout).



Exemplu (Corporate & HR Security):

Un algoritm analizează metadata: un angajat care de obicei pleacă la 17:00 începe să intre în clădire la ore atipice, descarcă volume mari de date sau accesează zone fizice unde nu are treabă, dar pentru care are drepturi (badge valid).

Sistemul ML calculează un "Risk Score".

Când scorul trece de 70, alertează preventiv Managerul de Securitate pentru o discuție, înainte ca incidentul să aibă loc.

Analiza de Risc/Nota Prudentă:

Etica și “Minority Report”

A acuza un om pe baza unei probabilități statistice este periculos și imoral.

ML trebuie să fie un indicator pentru investigație, nu judecător.



3

Pricing and Revenue Optimization (Optimizarea Bugetului și a Costurilor)

Securitatea este mereu sub presiune bugetară. Departamentele de securitate operează într-un paradox instituțional persistent: sunt evaluate prin absența evenimentelor negative, nu prin prezența lor.

Când nimic nu se întâmplă, managementul financiar vede cheltuieli greu de justificat, ore de pază, abonamente la sisteme de monitorizare, mentenanță preventivă, licențe software. Întrebarea ***"de ce cheltuim atât dacă nu s-a întâmplat nimic?"*** ignoră tocmai faptul că investiția în securitate este cauza directă a absenței incidentelor. Această invisibilitate a succesului face ca bugetele de securitate să fie primele tăiate în perioade de austeritate, creând vulnerabilități silențioase care se materializează dramatic tocmai când resursele sunt reduse.

ML ne ajută să justificăm fiecare leu cheltuit prin Analiză Cost-Beneficiu automată. Machine Learning transformă securitatea dintr-un centru de cost opac într-o funcție organizațională măsurabilă și optimizabilă.

Prin **corelarea continuă a datelor operaționale**, incidente produse, zone de risc, ore de patrulare, costuri de mentenanță, valoarea bunurilor protejate, **sistemele ML pot genera automat rapoarte de eficiență** care demonstrează concret randamentul fiecărei cheltuieli. Nu mai este nevoie de argumente intuitive în fața CFO-ului: datele arată obiectiv unde investiția produce securitate reală și unde consumă resurse fără impact proporțional. Această transparență analitică schimbă fundamental conversația bugetară – din negociere subiectivă în decizie bazată pe evidențe.



Aplicabilitate în Securitate:

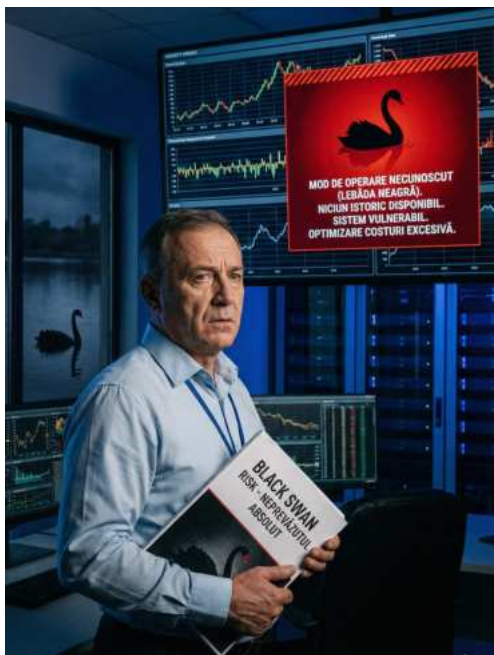
Optimizarea rutelor de patrulare și a mentenanței. A cheltui bugetul acolo unde este riscul real, nu unde "credem" noi că este.

Exemplu (Logistică și Pază):

Analizând istoricul de incidente și valoarea bunurilor din depozit, **ML sugerează**: "Ruta de patrulare actuală consumă 40% din timp în Zona A (unde nu au fost incidente de 3 ani). Mutați resursele în Zona C (unde sunt produse electronice scumpe)".

Rezultatul: Reduci numărul de ore de pază contractate cu 15%, păstrând același nivel de securitate (SLA), doar prin realocare inteligentă.





Analiza de Risc/Nota Prudentă:

Riscul “Lebedei Negre”

ML se bazează pe trecut. Nu poate prezice evenimente complet noi, nemaîntâlnite (ex: un nou mod de operare al infractorilor sau un atac terorist inedit).

Optimizarea excesivă a costurilor ne poate lăsa descoperiți în fața neprevăzutului absolut.

4

Predictive Maintenance (Mentenanța Predictivă)

Pentru inginerii și proiectanții de sisteme, acesta este cel mai valoros punct.





Aplicabilitate în Securitate:

Monitorizarea "sănătății" hard-disk-urilor din NVR-uri, a bateriilor din UPS-uri și/sau a mecanismelor de la turnicheți.



Exemplu (Sisteme Integrate):

În loc să schimbi bateriile sistemului de alarmă la fix 2 ani (preventiv - scump) sau când mor (*reactiv - riscant*), ML analizează curba de descărcare la fiecare test automat.

Sistemul îți spune: "*Bateria de la centrala 4 va ceda în 3 săptămâni*".

O schimbi exact la timp (Just-in-Time), eliminând downtime-ul și risipa.

Analiza de Risc/Nota Prudentă:

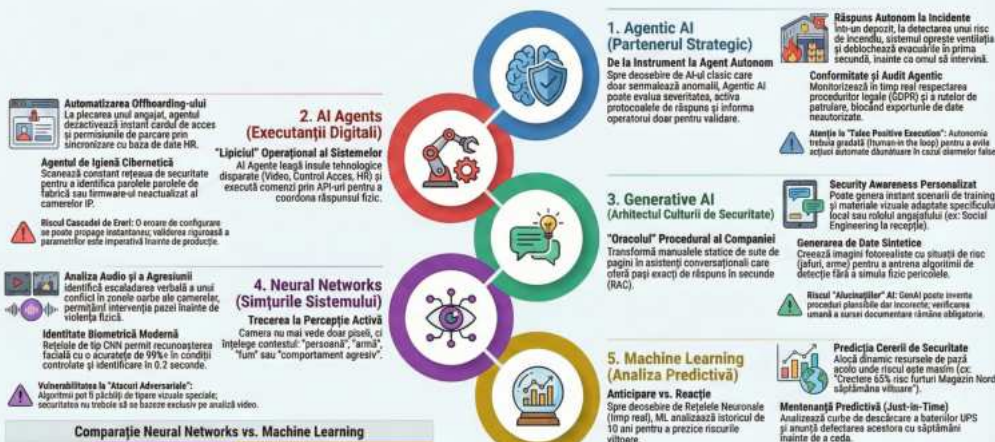
Dependența de senzori

Dacă senzorul care monitorizează bateria este defect, modelul predictiv este inutil.

Inspecția vizuală umană periodică rămâne o necesitate, chiar și în era AI.



Inteligența Artificială în Securitatea Fizică: Cele 5 Trepte ale Transformării



Comparație Neural Networks vs. Machine Learning

Caracteristică	Neural Networks	Machine Learning
Operare	Timp real, prezent	Analiză istorică + predicție
Funcție	"Ce se întâmplă ACUM?"	"Ce este PROBABIL să se întâmple?"
Ieșire (Output)	Alertă imediată (ex: intruz)	Recomandare strategică (ex: Dublare pază)

CURSURI DE SPECIALIZARE

Organizate de RQM Certification

www.rqmcert.ro

- ✓ **Consultant de securitate** - Ministerul Muncii
- ✓ **Manager de securitate** - Ministerul Muncii
- ✓ **Evaluator de risc la securitatea fizică** - Ministerul Muncii
- ✓ **Proiectant sisteme de securitate** - Ministerul Muncii
- ✓ **Inginer sisteme de securitate** - Ministerul Muncii
- ✓ **Tehnician sisteme de detecție, supraveghere video, control acces** - Ministerul Muncii
- ✓ **Managementul operațiunilor de securitate** - RQM Cert
- ✓ **Auditor de securitate cibernetică** - DNSC
- ✓ **Membru echipă CSIRT** - DNSC
- ✓ **Certified Ethical Hacker CIEH v13 AI** - EC Council
- ✓ **Certified Ethical Hacker CEH v13 iLearn** - EC Council
- ✓ **Certified Chief Information Officer (CCISO) iLearn** - EC Council
- ✓ **Certified Penetration Tester (CPENT) AI V2 iLearn** - EC Council
- ✓ **Certified SOC Analyst (CSA)** - EC Council
- ✓ **Certified Threat Intelligence Analyst (CTIA)** - EC Council
- ✓ **CISA – Certified Information Systems Auditor (Training)** - ISACA
- ✓ **CISM – Certified Information Security Manager (Training)** - ISACA
- ✓ **CRISC – Certified in Risk and Information Systems Control (Training)** - ISACA
- ✓ **Artificial Intelligence (AI) Foundation™** - Cloud Credential Council (CCC)
- ✓ **ISO 42001 Artificial Intelligence Management System** - PECB
- ✓ **Certified Artificial Intelligence Professional (CAIP)** - PECB

BIBLIOGRAFIE

Practices for Governing Agentic AI System

Sursa: OpenAI (2024/2025)

Link: <https://cdn.openai.com/papers/practices-for-governing-agentic-ai-systems.pdf>

Această lucrare este fundamentală pentru înțelegerea riscurilor asociate autonomiei. Documentul definește practici de siguranță vitale pentru implementarea Agentic AI descris în White Paper-ul analizat. Introduce concepte precum "limitarea spațiului de acțiune" (constraining action-space) esențială pentru a preveni ca un sistem de securitate să blocheze uși în mod eronat, și necesitatea unor mecanisme de "întrerupere" (kill switches). Acest studiu oferă baza teoretică pentru a elabora proceduri de *over-sight* uman, asigurând că "partenerul" digital nu devine o amenințare.

Agentic AI Security: Threats, Defenses, Evaluation, and Open Challenges

Sursa: arXiv (Cornell University), Datta et al. (Octombrie 2025)

Link: <https://arxiv.org/html/2510.23883v1>

Acest articol oferă o perspectivă pragmatică asupra modului în care Agentic AI este deja implementat în SOC-uri (Security Operations Centers). Descrie tranziția de la "detectarea mișcării" la "detectarea amenințării contextuale". Pentru cititorii interesați de ROI (Return on Investment), sursa explică cum agenții autonomi pot rezolva problema deficitului de personal, permițând monitorizarea a mii de camere fără a crește numărul de operatori umani, validând astfel pilonul "Operațiuni de Securitate Autonome" din ghid.

Identity Management for Agentic AI

Sursa: OpenID Foundation (Octombrie 2025)

Link: <https://openid.net/wp-content/uploads/2025/10/Identity-Management-for-Agentic-AI.pdf>

Deoarece agenții AI execută acțiuni în numele organizației (deschid uși, trimit notificări), ei necesită identități digitale securizate și auditate. Acest white paper tehnic abordează problema autentificării și autorizării agenților non-umani. Oferă arhitecților de securitate ghidaj despre cum să implementeze standarde precum **OAuth 2.1** pentru a asigura că agenții au doar permisiunile necesare (*Least Privilege*), prevenind riscul de "Zombie Accounts" menționat de mine în ghid.

NIST SP 800-215: Guide to a Secure Enterprise Network Landscape

Sursa: NIST (National Institute of Standards and Technology)

Link: <https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.800-215.pdf>

Deși este un standard general de securitate, NIST SP 800-215 este esențial pentru implementarea capacității de "monitorizare operațională" descrisă în document. Ghidul detaliază segmentarea rețelei și utilizarea analizei comportamentale (UEBA) pentru a detecta dispozitive compromise. Este o referință obligatorie pentru a transforma ideea de "agent de igienă" într-o politică de securitate conformă cu standardele internaționale.

What Is Retrieval-Augmented Generation (RAG)?

Sursa: NVIDIA / AWS / Palo Alto Networks

Link: <https://blogs.nvidia.com/blog/what-is-retrieval-augmented-generation/>

Acest ghid explică mecanismul tehnic din spatele "Oracolului Procedural": indexarea documentelor, transformarea în vectori și recuperarea contextului relevant pentru a preveni "halucinațiile" (inventarea de fapte). Pentru managerii de securitate, aceste resurse explică cum RAG asigură acuratețea răspunsurilor, un element critic atunci când deciziile bazate pe AI implică siguranța fizică.

Mitigating Security Risks in Retrieval Augmented Generation (RAG) LLM Applications

Sursa: Cloud Security Alliance (CSA)

Link: <https://cloudsecurityalliance.org/blog/2023/11/22/mitigating-security-risks-in-retrieval-augmented-generation-rag-llm-applications>

Această analiză detaliată abordează riscurile de "data poisoning" (otrăvirea bazei de cunoștințe) și injecția de prompturi. Oferă strategii defensive concrete, cum ar fi controlul accesului la nivel de document în cadrul vectorilor și sanitizarea inputurilor. Este esențială pentru a mitiga riscul de Data Leakage menționat în ghid, asigurând că asistentul AI nu devine un vector de atac intern.

Deep Learning for Video Surveillance: A Review

Sursa: JETIR / IEEE Xplore Context (2024)

Link: <https://www.jetir.org/view?paper=JETIR2409164>

Acest studiu academic oferă o privire "sub capotă" a algoritmilor (precum YOLO - You Only Look Once) utilizați pentru detecția în timp real. Explică diferența dintre analiza video bazată pe pixeli (veche, predispusă la alarme false) și cea bazată pe obiecte/semantică. Pentru cititorii tehnici, sursa validează capacitatea rețelelor neuronale de a reduce zgomotul și de a identifica amenințări complexe, precum armele ascunse sau comportamentele precursora violenței.

Guidelines 3/2019 on processing of personal data through video devices

Sursa: European Data Protection Board (EDPB)

Link:

https://www.edpb.europa.eu/sites/default/files/files/file1/edpb_guidelines_201903_video_devices_en_0.pdf

Aceasta este resursa de referință absolută pentru conformitatea legală. Ghidul clarifică distincția dintre supravegherea video simplă și prelucrarea datelor biometrice (care necesită un temei legal mult mai strict). Explică cerințele de informare, termenele de retenție și principiul minimizării datelor. Orice implementare a capabilităților biometrice descrise în White Paper trebuie să fie auditată prin prisma acestui document pentru a evita sancțiunile drastice.

Guidelines on the use of facial recognition technology in the area of law enforcement

Sursa: EDPB (Mai 2023/2024)

Link: **https://www.edpb.europa.eu/system/files/2023-05/edpb_guidelines_202304_frtlawenforcement_v2_en.pdf**

Deși orientat către autorități, acest ghid stabilește standardul de de facto și pentru securitatea privată de nivel înalt (infrastructură critică). Discută despre riscurile de discriminare algoritmică (bias) și necesitatea supravegherii umane stricte în cazul alertelor pozitive, susținând avertismentele etice din documentul analizat.

ISO 22341:2021 Security and resilience — Protective security — Guidelines for crime prevention through environmental design

Sursa: ISO (International Organization for Standardization)

Link: Pagina oficială a standardului

Acesta este primul standard internațional care codifică principiile CPTED. Pentru cititorii interesați de designul de securitate, standardul oferă cadrul metodologic pentru a integra elementele naturale de supraveghere cu tehnologiile AI descrise în document. Este "manualul" pentru arhitecți și consultanți de securitate care doresc să prevină riscurile încă din faza de proiectare a spațiului.

ISO/IEC 42001:2023 Information technology — Artificial intelligence — Management system

Sursa: ISO (International Organization for Standardization)

Link: Pagina oficială a standardului

Publicat recent, ISO 42001 este standardul de aur pentru guvernanta AI. Oferă cadrul pentru gestionarea riscurilor etice, operaționale și de securitate asociate sistemelor AI. Adoptarea acestui standard este recomandarea strategică supremă pentru orice organizație care implementează tehnologiile din White Paper, asigurând conformitatea cu viitorul EU AI Act și construind încredere în sistemele autonome.



www.rqmcert.ro



office@rqmcert.com



+40 356 173 020

RQM Certification

RQM Certification cu sediul în Timișoara este un furnizor de formare profesională cu o echipă excepțională de specialiști cu mare experiență în formare profesională, servicii de evaluare și audit. Compania are expertiză în domeniul sistemelor de management al calității, al mediului, al sănătății și securității la locul de muncă, al automobilelor, al securității fizice, al informațiilor și al serviciilor IT. Programele de formare sunt concepute pentru a sprijini învățarea activă în conformitate cu standardele internaționale și cerințele specifice fiecărei industrii.