



Civila tekniktrender med relevans för cyberförsvar

Explorativ analys av trender 2026

John Ziegenbein, Henrik Karlzén, Viktor Bergström,
Christoffer Lauri

John Ziegenbein, Henrik Karlzén, Viktor Bergström,
Christoffer Lauri

Civila tekniktrender med relevans för cyberförsvar

Explorativ analys av trender 2026

Titel	Civila tekniktrender med relevans för cyberförsvar – Explorativ analys av trender 2026
Title	Civilian Technology Trends with Relevance for Cyber Defence – Explorative Analysis of Trends 2026
Rapportnr/Report no	FOI-R--5956--SE
Månad/Month	Juni
Utgivningsår/Year	2026
Antal sidor/Pages	78
ISSN	1650-1942
Uppdragsgivare/Client	Försvarsmakten
Forskningsområde	Cyberförsvar och cybersäkerhet
FoT-område	Operationer i cyberdomänen
Projektnr/Project no	E38582
Godkänd av/Approved by	Foteini Papiri
Ansvarig avdelning	Cyberförsvar och ledningsteknik

Bild/Cover: Shutterstock, Martin Prochazkacz

Detta verk är skyddat enligt lagen (1960:729) om upphovsrätt till litterära och konstnärliga verk, vilket bl.a. innebär att citering är tillåten i enlighet med vad som anges i 22 § i nämnd lag. För att använda verket på ett sätt som inte medges direkt av svensk lag krävs särskild överenskommelse.

This work is protected by the Swedish Act on Copyright in Literary and Artistic Works (1960:729). Citation is permitted in accordance with article 22 in said act. Any form of use that goes beyond what is permitted by Swedish copyright law, requires the written permission of FOI.

Sammanfattning

Rapporten analyserar de tekniktrender som fem inflytelserika analysföretag identifierat som särskilt relevanta för 2026. Rapportförfattarna har bedömt trenderna utifrån relevansen för Försvarsmakten och specifikt för cyberförsvaret. Nio trender tas upp i rapporten:

- multiagenta system
- tillgång till AI-modeller och AI-beräkningskapacitet
- molntjänsters utmaningar
- AI-baserad mjukvaruutveckling
- AI för cyberförsvaret
- digitalt ursprung
- postkvantkryptering
- säkerhet för stora språkmodeller
- fysisk AI.

Varje trend analyseras utifrån forskning gjord av FOI och andra. Analyserna mynnar ut i praktiska rekommendationer till cyberförsvaret. Eftersom de flesta trender berör AI har flera av rekommendationerna gemensamma drag: skydd och försvar mot och med AI, organisatoriska förändringar samt aspekter kring digital suveränitet. Bland organisatoriska förändringar ingår även rekommendationer som rör postkvantkryptering.

Nyckelord: cyberförsvaret, AI, moln, suveränitet, postkvantkryptering

Summary

The report analyses the technological trends that five influential analysis firms have identified as particularly relevant for 2026. The authors of this report have assessed these trends based on the relevance to the Swedish Armed Forces, and specifically to cyber defence. Nine trends are addressed in the report:

- Multi-agent systems
- Access to AI models and AI computation
- Cloud service challenges
- AI-based software development
- AI for cyber defence
- Digital provenance
- Post-quantum cryptography
- Security for large language models
- Physical AI

Each trend is analysed based on research from FOI and other researchers. The analyses result in practical recommendations for the cyber defence. Since most trends involve AI, several of the recommendations share common themes: protection and defence against and with AI, organisational changes, and aspects of digital sovereignty. Among the organisational changes are also recommendations regarding post-quantum cryptography.

Keywords: cyber defence, AI, cloud, sovereignty, post-quantum cryptography

Innehållsförteckning

1	Inledning	7
1.1	Syfte och mål	7
1.2	Avgränsning	8
2	Bakgrund	9
2.1	Tidigare FOI-forskning om tekniktrender	9
2.2	Begrepp som rör AI, moln och kryptering	9
2.3	Populära företag och produkter för AI och moln	12
2.4	Försvarsmaktens strategier för AI och moln	14
2.5	Analysföretag som rapporterar om trender	15
3	Metod	16
3.1	Identifiering av trender	16
3.2	Exkludering av irrelevanta trender	17
3.3	Gruppering av inkluderade trender	17
3.4	Beskrivning och analys av grupperade trender	18
4	Tekniktrender för 2026	19
4.1	Multiagenta system	20
4.2	Tillgång till AI-modeller och AI-beräkningskapacitet	23
4.3	Molntjänsters utmaningar	27
4.4	AI-baserad mjukvaruutveckling	30
4.5	AI för cyberförsvar	34
4.6	Digitalt ursprung	38
4.7	Postkvantkryptering	41
4.8	Säkerhet för stora språkmodeller	45
4.9	Fysisk AI	49
5	Diskussion	53
5.1	Rekommendationer för att hantera trenderna	53
5.2	Förhållandet till Försvarsmaktens strategier	56
5.3	Kommer trenderna fortsätta?	57
5.4	Är dessa trender verkligen de mest relevanta?	58

6	Slutsatser	59
6.1	Vilka trender beskrivna av analysföretag är relevanta för cyberförsvaret?	59
6.2	Vad rekommenderas cyberförsvaret göra för att hantera trenderna?.....	60
6.3	Beskriver trenderna mogen teknik som kommer fortsätta vara relevant?	61
7	Referenser	62
Bilaga A	Extraherade trender	76

1 Inledning

Rapporten analyserar de tekniktrender som fem analysföretag identifierat som särskilt relevanta för 2026. Rapportens fokus är på trendernas påverkan på cyberförsvaret och analysen mynnar ut i rekommendationer till Försvarsmakten.

1.1 Syfte och mål

Rapporten är framtagen i ett projekt som är en del av Försvarsmaktens FoT-beställning 2026–2028. Projektet heter *Omvärldsanalys av civila koncept och teknik för cyberförsvaret*. Cyberförsvaret är skilt från cybersäkerhet, där det senare i huvudsak rör förebyggande skydd och finns i hela totalförsvaret, medan cyberförsvaret är mer aktivt agerande och en del av Försvarsmakten. Det mesta i rapporten rör främst cyberförsvaret, men det finns också delar (framförallt om moln samt kryptering) som är relevanta för andra delar av Försvarsmakten och övriga totalförsvaret.

Syftet med rapporten är att bidra med en del av svaren på projektets följande fem frågeställningar:

- Vilken kunskap om nya koncept och tekniker från civil forskning inom cybersäkerhet och cyberförsvaret är relevant att överföra till Försvarsmaktens cyberförsvaret?
- Vilka möjligheter och hot för cyberförsvaret ger den nya forskningen?
- Förväntas konceptet/tekniken ersätta nuvarande koncept/teknik hos leverantörer?
- Hur mogna är koncepten och teknikerna?
- Vilka frågor bör ställas för att avgöra mognadsgraden över tid hos AI-teknik som kan ha relevans för cyberförsvaret?

Målet med rapporten är att ge en översikt över civila trender med relevans för cyberförsvaret och att ge cyberförsvaret praktiska rekommendationer. Detta ska också kunna ligga till grund för mer omfattande bedömningar senare i projektet eller i andra forskningsprojekt. För detta kan projektets frågeställningar konkretiseras till rapportens frågor.

Rapporten besvarar följande frågor:

1. Vilka trender beskrivna av analysföretag är relevanta för cyberförsvaret?
2. Vad rekommenderas cyberförsvaret göra för att hantera trenderna?
3. Beskriver trenderna mogn teknik som kommer fortsätta vara relevant?

1.2 Avgränsning

Rapporten gör inte anspråk på att identifiera alla relevanta trender för Försvarsmaktens cyberförsvar, utan fokuserar på ett antal trender utifrån en bedömning att dessa har särskild relevans för cyberförsvar. Rapporten undersöker varje trend på ett fåtal sidor vardera. Detta bör ställas mot att trenderna kan representera egna omfattande forskningsområden. Rapporten ger därför bara en överblick av de ingående trenderna och kan ses som en uppsättning förstudier.

2 Bakgrund

Detta kapitel beskriver tidigare FOI-forskning om tekniktrender, övergripande begrepp som läsaren behöver vara bekant med, relevanta populära företag och deras produkter, Försvarmaktens strategier för AI och moln samt analysföretag som rapporterar om trender.

2.1 Tidigare FOI-forskning om tekniktrender

FOI gör ofta olika typer av omvärldsanalyser. Ett exempel från 2013 är memot Teknisk prognos cybersäkerhet [1]. Där redovisas vilken påverkan den generella it-utvecklingen bedömdes få på cybersäkerheten. Bland annat bedömdes cyberattacker ge en unik möjlighet att tillfälligt stänga av infrastruktur samtidigt som trenden gick mot att nästan vem som helst kunde ”sjösätta attacker, utan att inneha tekniska specialkunskaper”. Det beskrivs också att det funnits för stor tilltro till att marknaden skulle ordna säkra mjukvaror, utan större incitament. Ett annat exempel på omvärldsanalys är memot Urval av avskanningsämnen 2022 [2], där det identifieras intressanta ämnen att forska mer om. Bland annat nämns kvantkommunikation, digital teknik för hälsa, agil utvecklingsfilosofi samt ledning av cyberoperationer. Ett tredje exempel är en antologi om militärteknik från 2026 [3] där FOI-forskare från olika forskningsområden förutsäger de närmaste 25 årens utveckling inom respektive område, däribland cyberdomänen. Bland annat nämns en stundande kvantteknikrevolution, med kvantdatorer som om tjugo år tros kunna knäcka vissa av dagens krypton.

Det finns också exempel med mer praktisk koppling till Försvarmakten. Ett exempel är en form av omvärldsanalys som benämns perspektivstudier. I sådana studier bidrar FOI med sin kompetens för att ge stöd till överbefälhavarens militära råd inför kommande försvarsbeslut. Perspektivstudierna är väldigt långsiktiga. Hur FOI arbetat med Försvarmakten i dessa studier beskrivs bland annat i rapporten [4]. Därtill har FOI ett forskningsprogram om svaga signaler, där det bedrivs datadriven detektion och analys av tidiga tecken på ny relevant forskning med betydelse för totalförsvaret [5].

2.2 Begrepp som rör AI, moln och kryptering

De följande tre tabellerna förklarar övergripande begrepp som läsaren behöver vara bekant med, inom AI, moln och kryptering.

Tabell 1 förklarar de övergripande begrepp om artificiell intelligens (AI) som läsaren behöver vara bekant med.

Tabell 1. Förklaringar av begrepp som rör AI.

Begrepp	Förklaring
Artificiell intelligens (AI)	Ett paraplybegrepp för tekniker eller verktyg där datorer utför aktiviteter som tidigare krävt mänsklig intelligens. Aktiviteterna kan bland annat röra lärande, perception, problemlösning och beslutsfattande. Används bland annat för bildigenkänning, spel och textgenerering. Ofta kräver teknikerna bakom stora mängder data och beräkningskraft (en dators förmåga att utföra beräkningar).
AI-modell	En typ av AI som byggts upp med träningsdata och som kan göra prediktioner eller ta beslut. Kan också kallas <i>en AI</i> .
Generativ AI	En sorts AI-teknik som lär sig av träningsdata och sedan skapar (genererar) nytt innehåll i form av exempelvis text, bilder, video eller ljud. Genereringen styrs typiskt sett av en instruktion från användaren (en prompt) via en chattfunktion. Numera är generativ AI så dominerande att det ofta används synonymt med AI-modell och AI.
Stor språkmodell	En variant av generativ AI, tränad på stora mängder text för att tolka och generera mänskligt språk. De flesta välkända AI-modeller är stora språkmodeller och de kan oftast interageras med via en chattfunktion. På engelska är begreppet <i>large language model</i> (förkortat LLM).
AI-agent	En mer autonom AI som kan agera målinriktat och självständigt, snarare än bara svara på frågor.
Träningsdata	De data som ges som indata till AI-modellen för att den ska hitta mönster. Exempelvis kan en AI-modell tränas på finansiella dokument varefter den lärt sig om finansportföljer och investeringsstrategier.
Källkod	Källkod är människoläsbara instruktioner till en dator och dessa översätts sedan till maskinkod (ettor och nollor).
Prompt	Instruktioner eller indata till en AI-modell. Instruktionerna kan ges i förväg för att vara relevanta under längre tid, eller ges vid varje användning av modellen och då exempelvis i chattfönster.
Modellvikter	En AI-modells inlärda parametrar och konfigurationer. Det finns öppna och stängda (proprietära) modellvikter. Öppna modellvikter innebär att leverantören har offentligt tillgängliggjort modellens inlärda parametrar och konfigurationer. Detta möjliggör att driftsätta AI-modellen på lokal hårdvara. Detta innebär även att det är möjligt att finjustera modellen med egna träningsdata. Stängda modellvikter innebär att modellen enbart finns tillhandahållen via leverantörers gränssnitt.

Begrepp	Förklaring
Öppen källkod och öppen träningsdata	Utöver modellvikter kan även träningsdata och källkod vara öppen. Det innebär alltså att leverantören offentligt tillgängliggjort modellens källkod och träningsdata.
Hallucination	Felaktigt genererade utdata från en AI-modell, exempelvis sådant som AI-modellen uppger vara fakta men som är falskt.
Sele	Någon form av infrastruktur runt AI-modellen. Blev ett populärt begrepp i början av 2026 och det finns än så länge lite olika tolkningar av begreppet. Kan exempelvis anses innehålla orkestrering, andra verktyg och skyddsräcken [6]. Har inget vedertaget namn på svenska. På engelska är begreppet <i>harness</i> , vilket även används som verb.
Skyddsräcke	En samling säkerhetsmekanismer för att begränsa AI-modellers beteende till vad som anses säkert. Exempel på skyddsräcken är in- och utmatningsfilter samt åtkomstkontroll. Kallas även skyddsspärr på svenska. På engelska är begreppet <i>guardrail</i> .

Tabell 2 förklarar de övergripande begrepp om moln som läsaren behöver vara bekant med.

Tabell 2. Förklaringar av begrepp som rör moln.

Begrepp	Förklaring
Moln	Moln (eller molntjänster) erbjuder tillgång till datorkapacitet från delade fysiska resurser. Fler resurser kan snabbt tilldelas, med minimal interaktion mellan tjänsteleverantör och användare.
Hyperscaler	De globalt största leverantörerna av traditionella moln.
Multimoln	En kombination av flera separata moln från olika leverantörer.
Neocloud	En typ av moln avsedda att användas för AI, med särskild AI-anpassad hårdvara.
Konfidentiell databehandling	Ett koncept där data skyddas under bearbetning på en annan aktörs hårdvara.
Geopatriering	När användaren flyttar sina data från ett moln där användare har mindre kontroll (ofta hos en hyperscaler) till infrastruktur där användaren har mer kontroll.

Tabell 3 förklarar de övergripande begrepp om kryptering som läsaren behöver vara bekant med.

Tabell 3. Förklaringar av begrepp som rör kryptering.

Begrepp	Förklaring
Postkvantkryptering (PQC)	Kryptografiska algoritmer som är avsedda att vara säkra även om kvantdatorer realiserar.
Kvantdator	En typ av dator som förväntas vara överlägsen klassiska datorer för vissa typer av beräkningar. Kvantdatorerna utnyttjar att kvantmekaniken tillåter att datorn håller en sorts kombination av flera värden (1 och 0) samtidigt. Än så länge finns bara mycket begränsade kvantdatorer.
Symmetriska krypton	En typ av krypto där samma nyckel används för både kryptering och dekryptering.
Asymmetriska krypton	En typ av krypto där olika nycklar används för kryptering och dekryptering.
Heltalsfaktorisering	Att dela ett heltal i flera faktorer (tal), där faktorerna oftast ska vara primtal. Exempelvis kan 6 delas upp i faktorerna 2 och 3, eftersom $6 = 2 * 3$.

2.3 Populära företag och produkter för AI och moln

Detta avsnitt beskriver kortfattat de största företagen och produkterna för AI och moln. Företag och produkter för kryptering beskrivs inte eftersom relevansen för kryptering i rapporten främst är för postkvantkryptering, där fokus är på algoritmer snarare än på produkter. Därtill är de produkter för postkvantkryptering som än så länge finns typiskt generella webbläsare och andra mjukvaror med bredare tillämpningsområde än kryptering.

Tabell 4 beskriver kortfattat de största AI-företagen och deras namnkunniga modeller enligt Garrido-Merchán m.fl. [7].

Tabell 4. Populära AI-modellserier 2026. Versioner av modeller har förenklats till modellserier, medan vidareutvecklingar exkluderats.

Land	Företag	Modellserie	Öppna modellvikter
USA	Anthropic	Claude	Nej
	Google	Gemini	Nej
		Gemma	Ja
	XAI	Grok	Nej
	Open AI	GPT	Nej
		GPT-oss	Ja
	Meta	Llama	Ja
Kina	Microsoft	Phi	Ja
	Zhipu AI	GLM	Ja
	Moonshot AI	Kimi	Ja
	Deepseek	Deepseek	Ja
	Minimax	Minimax	Ja
	Alibaba	Qwen	Ja
	Step Fun	Step	Ja
	Xiaomi	Mimo	Ja
	Ant group	Ling	Ja
	Tencent	Hunyuan	Ja
	Shanghai AI Lab	Intern LM	Ja
	01.AI	Yi	Ja
	Baichuan	Baichuan	Ja
	Frankrike	Mistral AI	Mistral

Dessa AI-modeller kan också integreras med andra verktyg och byggas vidare på (exempelvis i AI-selar). Ett exempel på detta är Microsoft Copilot som bygger på Open AI GPT.

Tabell 5 beskriver de största molnleverantörerna och deras moln enligt [8]. Därtill beskrivs deras marknadsandelar.

Tabell 5. Populära moln 2026.

Land	Företag	Tjänst	Marknadsandel
USA	Amazon	AWS	30 %
	Microsoft	Azure	20 %
	Google	Cloud Platform	13 %
	Oracle	Cloud	4 %
	Salesforce	Cloud	3 %
	IBM	Cloud	2 %
Kina	Alibaba	Cloud	2 %
	Tencent	Cloud	2 %

2.4 Försvarsmaktens strategier för AI och moln

Under 2025 publicerade Försvarsmakten både en strategi för AI [9] och en strategi för hybrida multimoln [10].

I *Försvarsmaktens strategi för artificiell intelligens (AI)* beskrivs att ”slagfältet formas av kampen om informationsöverläge, förmåga till anpassning i realtid och kortare beslutscykler”. Strategin beskriver att AI skapar ”förutsättningar för att konsekvent ligga steget före en motståndare”. Strategin beskriver också att AI ”är en teknik med potential att omforma stora delar av Försvarsmaktens verksamhet” och att införandet därför behöver förstås som ”en strategisk transformation, inte som en isolerad utvecklingsfråga”. Strategin har en planeringshorisont för 2025–2030 och nämner ”insikten att teknikens snabba utveckling gör längre prognoser osäkra”. Därtill beskriver strategin att ”AI-system ska utformas och införas med höga krav på säkerhet och resiliens”.

I *Försvarsmaktens hybrida multimolnstrategi* beskrivs ”användning av molntjänster från flera molnleverantörer i kombination med privata moln”. Sådana lösningar kan minska problemen som uppstår om ett moln stängs ned av leverantören eller på grund av vissa typer av driftstörningar. Molnstrategins syfte är att ”ge en gemensam inriktning för hur Försvarsmakten ska nyttja molntechnologi som en strategisk resurs” samt att implementera Natos initiativ till gemensam molninfrastruktur i alliansen. Målbilden är att gå från ett fokus på system, domän och nät till att bli datacentrerat. Detta beskrivs som att mer data ”tillgänglig i rätt tid och med rätt kvalitet ökar både handlingskraft och förmågan att fatta beslut, vilket är avgörande för att nå operativ effekt”. I strategin beskrivs bland annat tre typer av suveränitet som berör om data lagras i egen jurisdiktion, om det tekniskt går att byta mellan moln samt om det finns egen driftpersonal.

2.5 Analysföretag som rapporterar om trender

Det finns många företag som på ett eller annat sätt publicerar analyser om trender. Det finns exempelvis analysföretag inom it och teknikanalys, såsom Gartner, Forrester och IDC. Det finns också analysföretag som är inriktade på revision och konsultverksamhet, däribland Deloitte, KPMG, EY och PWC. Därtill finns it-inriktade konsultbolag såsom Capgemini och Juniper Research (ej relaterat till Juniper Networks). Dessutom finns det managementkonsultbolag som McKinsey och Bain, vilka också i viss mån publicerar offentliga rapporter. Denna rapport baseras på trender identifierade av de tre första typerna av analysföretag, men inte på managementkonsultbolag. Mer specifikt använder rapporten källor från de analysföretag som beskrivs i tabell 6. Motiveringen till valet beskrivs i avsnitt 3.1.

Tabell 6. Inkluderade analysföretag.

Företag	Beskrivning
Gartner	Gartner är ett av de ledande analysföretagen inom it och teknikanalys [11].
Juniper Research	Juniper research marknadsför sig själva som specialister inom digital teknik för finans, telekom och sakernas internet [12]. Enligt en undersökning är de ett av de mest inflytelserika analysföretagen, tillsammans med bland andra Gartner och Forrester [13].
Capgemini	Capgemini är ett stort konsultbolag som specialiserar sig på rådgivning inom it. En av företagets styrkor sägs vara att utföra marknadsanalyser inom området [14].
Deloitte	Deloitte är ett av världens ledande tjänsteföretag och verkar som ett globalt nätverk av juridiskt oberoende medlemsfirmor [15]. Deloitte är ett av världens största revisionsföretag.
Forrester	Forrester är ett av de ledande analysföretagen inom it och teknikanalys [11].

3 Metod

Detta kapitel beskriver den metod som använts för att ta fram och beskriva tekniktrenderna. Kapitlet beskriver identifiering av trender, exkludering av irrelevanta trender, gruppering av inkluderade trender samt beskrivning av trenderna.

3.1 Identifiering av trender

För att identifiera tekniktrender användes källor i form av trendrapporter från analysföretag, där trendrapporterna matchar internetsökningar efter trendrapporter för 2026.

Den här rapporten nyttjar trendrapporterna just som en källa till identifierade tekniktrender samt i viss utsträckning för att ge uppslag till explorativ analys av trenderna. Som framgår i avsnitt 3.4 använder analysen även andra källor och då främst i form av FOI-rapporter, FOI-memon och forskningsartiklar från olika forskare.

De företag och trendrapporter som inkluderas återges i tabell 7. Forresters trendrapport inkluderas trots att dess titel använder ordet prediktion (eng. predictions). Detta motiveras delvis av att Juniper Research kallar sina trender framväxande (eng. emerging), vilket visar på likhet mellan trend och prediktion. Därtill mynnar denna rapports analys ut i framåtblickande potentiella konsekvenser och rekommendationer. Någon skillnad i ovisshetsbedömning görs inte mellan trender och prediktioner.

Tabell 7. Inkluderade trendrapporter, med översatta namn.

Analysföretag	Trendrapport
Gartner	Topp 10 strategiska tekniktrender för 2026 [16]
Juniper Research	Topp 10 framväxande tekniktrender 2026 [17]
Capgemini	Topp tekniktrender för 2026 [18]
Deloitte	Tekniktrender 2026 [19]
Forrester	Prediktioner 2026 – teknik och säkerhet [20]

Ett problem är att det saknas transparens kring vilka metoder företagen använder för att ta fram rapporter av det här slaget [21]. Företagen har dessutom fått kritik för att ha ovetenskapliga metoder [22]. I de här utvalda trendrapporterna är det bara Juniper Research som (mycket ytligt) beskriver sin metod. De säger sig ha utgått från sin kunskap om marknaden och sedan valt trender som de förväntar sig kommer ha störst påverkan på marknaden. Förutom den begränsade transparensen har flera analysföretag anklagats för att leverera rapporter med påhittade referenser, vilka hallucinerats fram av AI.

Exempelvis gäller sådana anklagelser EY, McKinsey och Deloitte [23] samt KPMG [24]. Deloitte är ett av analysföretagen som trots allt inkluderas som källa i denna rapport. Författarna till denna rapport har kontrollerat att källorna som uppges i trendrapporterna existerar. I de flesta trendrapporter finns högst begränsat med referenser. Just Deloitte har flest, men de går ganska ofta till egna undersökningar som inte är publikt tillgängliga. Inga tecken på hallucinerade referenser hittades.

3.2 Exkludering av irrelevanta trender

Analysföretagens trendrapporter identifierar totalt 37 trender. Bilaga A har en lista på alla trender och var de kommer från.

Fyra av trenderna exkluderades av denna rapports författare. I samtliga fall skedde exkluderingen eftersom källan inte behandlar trendens säkerhets- eller försvarsaspekter. I tabell 8 beskrivs vilket företag som tar upp trenden och hur företaget beskriver den.

Tabell 8. Exkluderade trender.

Företag	Beskrivning
Deloitte	Trendföretaget beskriver att AI är förknippat med åtta signaler på teknik som kan komma att bli viktiga trender tillsammans med AI. Dessa signaler beskrivs med ett stycke var i analysföretagets trendrapport.
Forrester	Trendföretaget beskriver att AI ibland införs utan plan och att företagsledare kommer kräva mer av teknikchefer på denna punkt framöver.
Forrester	Trendföretaget beskriver att AI ibland inköps utan plan och att företagsledare kommer kräva mer av ekonomichefer på denna punkt framöver.
Juniper Research	Trendföretaget beskriver att laddningsbehoven för elfordon kommer att lösas med trådlös laddning framöver.

3.3 Gruppering av inkluderade trender

De 33 trenderna grupperades i nio mer övergripande trender. Grupperingen skapades utifrån bedömningar av denna rapports författare. Dessa bedömningar baserades på gemensamma faktorer mellan de underliggande trenderna. De större grupperna var mest uppenbara och skapades först. Därefter blev några underliggande trender över som inte passade in och som därför behövde nya grupper. Två grupper fick till slut bara en medlem (underliggande trend) vardera. En underliggande trend (om AI för försvar och säkerhet) delades upp i två övergripande trender (en om försvar och en om säkerhet).

Alla utom en av grupperna (om postkvantkryptering, se avsnitt 4.7) handlar i någon mån om AI. Kortfattat handlar AI-grupperna om att AI kan verka tillsammans i multiagenta system (4.1), kan generera mjukvara (4.4), skapa desinformation (4.6), försvara och angripa (4.5) samt fysiskt påverka (4.9). Dessa möjligheter ställer krav på säkerheten hos AI (4.8), innebär utmaningar för hårdvara och suveränitet (4.2) samt innebär nyttjande av molntjänster (4.3). Molntjänsterna är även relevanta för traditionell datoranvändning som inte använder AI.

3.4 Beskrivning och analys av grupperade trender

Varje övergripande trend beskrivs i fem steg:

1. **Vad säger analysföretagen?**

Detta sammanfattar relevanta delar av analysföretagens texter med fokus på drivkrafter, nuläge och framtidsaspekter. Denna del avgör avgränsningarna för trenden.

2. **FOI-forskning**

Detta beskriver forskning som FOI har gjort inom området. I vissa fall är denna forskning väldigt bred och täcker stor del av vad som ingår i trenden. I andra fall är FOI-forskningen mer punktinsatser som endast beskriver en begränsad del av trenden.

3. **Explorativ analys**

Detta beskriver trenden utifrån några vetenskapliga källor samt i viss mån internationella myndighetsrapporter, nyhetsartiklar och liknande informella källor (via Google). När det finns gott om tidigare FOI-forskning om trenden baseras analysen främst på den forskningen.

4. **Trendens konsekvenser för cyberförsvar**

Detta inkluderar både konsekvenserna av Försvarsmaktens möjligheter och att en potentiell hotaktör nyttjar trenden. Detta baseras på de tidigare texterna om trenden tillsammans med dokument från Försvarsmakten samt rapportförfattarnas erfarenheter.

5. **Rekommendationer till Försvarsmakten**

Detta handlar om att ge konkreta råd för hantering av trenden och baseras i huvudsak på respektive text om konsekvenser samt rapportförfattarnas erfarenheter. Varje rekommendation har ett unikt ID.

4 Tekniktrender för 2026

I detta kapitel beskrivs de övergripande trender som är relevanta ur ett cyberförsvarsperspektiv, enligt bedömningar av denna rapportts författare. I tabell 9 ges en översikt över de inkluderade övergripande trenderna. De underliggande trenderna identifierades i analysföretags trendrapporter. Denna rapportts fortsatta explorativa analys utgår i huvudsak från FOI-rapporter, FOI-memon och forskningsartiklar. Trendrapporterna ger därmed enbart uppslag till analysen av trenderna.

Tabell 9. Övergripande trender från analysföretags publikationer för tekniktrender för 2026. Kolumnen # anger antal extraherade trender.

#	Övergripande trend	Förklaring	Avsnitt
7	Multiagenta system	Enskilda specialiserade AI-agenter har börjat koordineras för att tillsammans utföra arbetsflöden med flera steg och nå mer komplexa mål.	4.1
6	Tillgång till AI-modeller och AI-beräkningskapacitet	Allt mer AI införs vilket kraftigt ökar behovet av beräkningskapacitet. Organisationer får svårt att få lämplig tillgång till AI-modeller.	4.2
5	Molntjänsters utmaningar	Den säkerhetspolitiska utvecklingen ökar behoven av bättre kontroll över data och tjänster, vilket ställer krav på molntjänsterna.	4.3
5	Fysisk AI	Organisationer försöker överföra vinsterna med digital AI till den fysiska omgivningen.	4.9
3	AI-baserad mjukvaruutveckling	Organisationer ser produktivetsfördelar och lägre kostnader med att utveckla mjukvara med hjälp av AI.	4.4
2	AI för cyberförsvar	AI missbrukas allt mer av angripare, vilket också ökar intresset av proaktivt AI-försvar.	4.5
2	Postkvantkryptering	Organisationer har lagt mer fokus på att planera för byten av kryptoalgoritmer, för att motstå de nya angrepp som skulle realiseras med kvantdatorer om sådana blev bättre.	4.7

#	Övergripande trend	Förklaring	Avsnitt
2	Säkerhet för stora språkmodeller	Nya säkerhetsmekanismer införs som svar på de nya säkerhetsproblem som uppstår när AI-modellerna blir allt mer kraftfulla.	4.8
1	Digitalt ursprung	Organisationer står inför ökade risker med kodmanipulation och AI-driven desinformation, vilket ger behov av bättre verifiering hos tredjepartsmjukvara och AI-genererat innehåll.	4.6
33	(Summa)	–	–

4.1 Multiagenta system

Trenden multiagenta system handlar om att enskilda specialiserade AI-agenter börjat koordineras för att tillsammans utföra arbetsflöden med flera steg och nå mer komplexa mål. De specialiserade agenterna utgörs ofta av domänspecifika språkmodeller. Ett resulterande multiagent system kan ses som agentisk AI [25]. En central utmaning med multiagenta system är samordningen mellan agenter, vilket kan ställa krav på omfattande organisatoriska förändringar.

4.1.1 Vad säger analysföretagen?

Gartner framhåller att multiagenta system är på uppgång eftersom en enskild AI-agent har svårigheter att utföra processer eller arbetsflöden med flera steg. De tror också att utvecklingen kommer börja med multipla agenter på en och samma plattform, till agenter på olika plattformar som kommunicerar, till slutligen ett internet av agenter. Gartner nämner också domänspecifika språkmodeller (eng. domain-specific language models, DSLM) som ett exempel på specialisering av agenter. En DSLM är en AI-modell som tränats på specialiserade data med fokus på en viss uppgift eller domän. Gartner beskriver att DSLM:er kan minska felen, accelerera införandet och minska kostnaderna för kritiska arbetsflöden som finans, hälso- och sjukvård samt hr. Gartner tror att mer än hälften av alla generativa AI-modeller för företag kommer vara domänspecifika 2028.

Juniper Research ser en trend i att organisationer som inför AI-stöd även investerar i DSLM:er i strävan efter mer korrekthet och precision, exempelvis vad gäller regelefterlevnad.

Deloitte tror att fördelar med agentisk AI kräver att hela verksamheter designas om snarare än att existerande processer automatiseras. Deloitte menar också att agenter är en form av arbetskraft som utmärker sig när det kommer till

definierade processer, medan mänsklig arbetskraft fortsätter vara nödvändig för att hantera förändringar i verksamhetskrav och komplexa problemscenarion.

Capgemini menar att nästa steg inom AI kommer handla om övergripande arbete med arkitekturer och integration snarare än om specifika verktyg eller AI-modeller. Många organisationer kommer också fokusera på DSLM:er för att uppnå bättre kontroll över resultat och prestanda. Capgemini spår även att organisationer kommer börja flytta fokus från kostnadsbesparande åtgärder till värdeskapande aktiviteter i sina AI-initiativ. Därtill beskrivs att organisationer kommer behöva designa om hela verksamheter snarare än att försöka automatisera existerande processer.

4.1.2 FOI-forskning

Tabell 10 beskriver de rapporter och memon som FOI publicerat som diskuterar multiagenta system.

Tabell 10. FOI-forskning som tar upp de koncept för multiagenta system som lyfts av analysföretagens trendrapporter.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Utredning av stödverktyg och metodik för utvärdering av AI-metoder	FOI-R--5453--SE [26]	2023	Rapporten nämner kortfattat multiagenta system.
Large Language Models in Defence: Challenges and Opportunities	FOI-R--5544--SE [27]	2024	Rapporten behandlar träning (finjustering) av LLM:er för att anpassa dem till en viss domän. Detta beskrivs leda till märkbara förbättringar, utan allt för stora investeringskostnader. Den stora svårigheten är att veta hur modellens tillförlitlighet ska utvärderas.

4.1.3 Explorativ analys

Forskningen om multiagenta system är etablerad som forskningsområde sedan lång tid och har växt stadigt över det senaste dryga decenniet [28]. Inom cybersäkerhet kan multiagenta system tänkas tillföra ett stöd för att hantera exempelvis överbelastningsförsök, spoofing och för att utföra inträngsdetektion [29].

I multiagenta system används flera olika specialiserade agenter. Vid specialiseringen, exempelvis med domänspecifika språkmodeller, är en viktig förutsättning tillgång till representativa och relevanta träningsdata samt beräkningskapacitet för träningsmomentet. En svårighet med specialiseringen är att om modellen tränas för mycket på specifika data så tappar den viktig generell förmåga [27]. Relaterat till detta är också att modellerna kan få svårt att minnas tidiga data eller instruktioner när det blir mycket nya data att behandla [30].

I multiagenta system finns det risker med att flera modeller och verktyg förlitar sig på varandras resultat, vilket kan innebära koordineringsutmaningar och kaskadeffekter. En vägledning avseende möjliga risker och rekommendationer har tagits fram av cybermyndigheter från Australien, USA med flera [31].

För att kommunicera mellan agenter och med verktyg finns det protokoll som Agent Communication Protocol (ACP) för kommunikation mellan agenter samt Model Context Protocol (MCP) för verktyg [32].

4.1.4 Konsekvenser för cyberförsvar

Multiagenta system innebär mer AI på fler ställen och med en förhoppning om att summan är större än delarna. Detta kan göra det snabbare att utföra angrepp, eller möjliggöra mer komplexa sådana. Det finns också skäl att tro att försvarare kommer kunna dra nytta av multiagenta system. Tekniskt kan det också bli lättare att uppnå segmentering i systemet, jämfört med en monolitisk implementation. Dock skapas koordineringsutmaningar mellan agenterna. Det är också en öppen fråga ifall övergången från en AI till flera AI ökar effekten linjärt, multiplikativt eller inte alls. Möjligheten för negativa kaskadeffekter bör också beaktas, exempelvis där en AI hallucinerar och nästa AI bygger på den förstas resultat.

Det kan finnas ökat behov av verksamhets- och organisationsförändringar till följd av den stora mängden agenter och koordinering mellan dem. Det kan rentav behövas nya organisationsenheter samt överväganden om att se AI-agenter som medarbetare. Därtill behöver kontinuitet och lärande över tid säkerställas så att inte AI ersätter kompetens som senare behövs hos människor. Dessutom bör det göras avvägningar mellan generalitet och specificitet hos AI-modellernas kompetens. Domänspecifik AI kan visserligen vara specialiserad, men dess specialistträning kan också innebära bortfall av viktig generell information på grund av tekniska begränsningar med inlärningstekniken.

4.1.5 Rekommendationer till Försvarsmakten

I följande punktlista beskrivs rekommendationer till Försvarsmakten avseende multiagenta system:

- R1. Analysera hur multiagenta system kan koordineras utan att kontrollen av agenterna övergår till motståndare.
- R2. Analysera hur ett AI-baserat cyberförsvar skulle se ut organisatoriskt när det kompletteras av multiagenta system.

4.2 Tillgång till AI-modeller och AI-beräkningskapacitet

Trenden tillgång till AI-modeller och AI-beräkningskapacitet handlar om att allt mer AI införs vilket kraftigt ökar behovet av beräkningskapacitet. De stora kraven på beräkningskapacitet gör i sin tur att få aktörer kan bygga egna AI-modeller som kan mäta sig med de bästa. Detta ger en utmaning med att få tillgång till AI-modeller och på rätt sätt.

Att införskaffa modeller från leverantörer och köra dessa lokalt är en möjlighet, men sådana presterar typiskt sämre. Att istället köra modellerna externt ger dock problem med konfidentialitet och tillgänglighet. Oavsett lokal eller extern användning gör den bristande insynen i externt tränade AI-modeller att resultaten måste tolkas med försiktighet.

4.2.1 Vad säger analysföretagen?

Gartner ser en ökande efterfrågan på AI-datacenter (även kallade AI-superdatorer) på grund av det grundläggande behovet av datorkraft för att tillverka och driftsätta AI-modeller.

Juniper Research menar att AI kraftigt ökat efterfrågan på datorkraft och energikonsumtion, vilket kommer skapa påfrestningar på elnäten. De menar att kärnkraft i form av småskaliga modulära reaktorer (eng. *small modular reactors*) är en potentiell lösning och att dessa kommer få brett regulatoriskt stöd under 2026. Juniper Research tror även att energibehovet kan minska genom neuromorfa kretsar (eng. *neuromorphic chips*) där hårdvara anpassas för vissa AI-beräkningar. Samma källa nämner även att kylning är ett problem för dagens integrerade kretsar för AI. De förväntar sig att detta bemöts med mikrofluidik (eng. *microfluidics*) vilket beskrivs som en teknik för att manipulera små mängder vätska genom väldigt små kanaler. Detta beskrivs kunna användas inom kylsystem för att leda kylvätskan närmare systemen som behöver kylas.

Deloitte tror att företag framöver kommer att försöka nå mer kostnadseffektiv AI-användning. Exempelvis kommer företag överväga att köra AI-modeller

lokalt när kostnaderna för drift och införskaffning av modellerna är lägre än avgifterna vid användning av externa AI-modeller. Därtill menar Deloitte att suveränitetskrav är en viktig faktor. Exempelvis är lokala AI-modeller intressanta, där den enskilda organisationen har rådighet och det är lättare att upprätthålla konfidentialiteten. En annan aspekt som Deloitte nämner som relevant är latens, där tidskritisk användning kräver att AI förs närmare klienterna, eller rentav sker i dem.

Capgemini beskriver en ökande trend med suveräna AI-ekosystem. Trenden möter problemet med att förlita sig på externa tjänster där insyn saknas. Att ha insyn beskrivs som särskilt viktigt inom offentliga tjänster, finans, hälso- och sjukvård samt försvar.

4.2.2 FOI-forskning

Tabell 11 beskriver de rapporter och memon som FOI publicerat som diskuterar tillgång till AI-modeller och beräkningskapacitet.

Tabell 11. FOI-forskning som tar upp de koncept för AI-modeller och beräkningskapacitet som lyfts av analysföretagens trendrapporter.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Artificiell intelligens och militära operationer	Memo 7807 [33]	2022	Memot följer kunskapsuppbyggande verksamhet vid FOI under 2021. Memot ger en bakgrund till AI-området och närliggande koncept. Memot beskriver flera begrepp som även lyfts i analysföretagens trendrapporter, däribland neuromorfa kretsar, moln och beräkningar nära klienter.
Effektiva AI-beräkningar 2024	Memo 8840 [34]	2025	Memot är på en sida och sammanfattar en studie om effektiva AI-beräkningar. Neuromorfa kretsar tas upp.
Stora språkmodeller som ett stödverktyg vid militär taktisk planering – Slutsatser från två tillämpade studier av stora språkmodeller	FOI-R--5852--SE [35]	2026	Rapporten undersöker hur stora språkmodeller kan användas vid militär taktisk planering. Kortfattat nämns problem med konfidentialitet vid användandet av externa AI-modeller.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
DeepSeek: Ändamålsenlighet och tillförlitlighet	FOI-R--5787--SE [32]	2026	Rapporten utvärderar kinesiska språkmodellen Deepseek ur ett tillförlitlighets- och ändamålsperspektiv med fokus på användning inom svensk myndighetsutövning.

4.2.3 Explorativ analys

Användningen av AI-modeller har ökat snabbt och mycket. Ökningen gäller även modellernas storlek och därmed behovet av kraftfull hårdvara. Hårdvarubehoven är störst vid träning. Av de tjugo största AI-datacentren är fjorton i USA, fyra i Kina, ett i Tyskland och ett i Norge [36]. Storleken på datacentren varierar stort. Företaget Xai hävdar att deras datacenter *Colossus*, vilken används för att träna deras AI-modell *Grok*, är världens kraftfullaste dator för träning. Colossus har 200 000 GPU:er¹, 194 Petabyte/s i minnesbandbredd, nätverksbandbredd på 3,6 Terabit/s och 1 Exabyte i lagringskapacitet [37]. Det kan jämföras med Norges datacenter som har 16 400 GPU:er. Kapprustningen för AI påverkar även it-system i övrigt. Behovet av minne i AI-infrastruktur har ökat priserna för minnen med ungefär 80 % under första kvartalet 2026 [38].

AI-infrastrukturerna kräver stora mängder el och mängderna överstiger vad de regionala elnäten kan tillgodose [39]. Företag som skapar och förvaltar datacenter utforskar därför möjligheterna att använda kärnkraft för att förse datacentren med el genom småskaliga modulära reaktorer [40].

Att utveckla egna kraftfulla AI-modeller ställer därmed stora krav på resurser. Inom Sverige och EU finns visserligen olika samarbeten med syfte att ta fram AI-modeller, däribland samarbetena AI Sweden, Wasp och Openeuro LLM. Detta är även kopplat till suveränitetskrav som följer av regelverk. Än så länge dominerar dock USA i byggandet av infrastrukturer för träning av AI-modeller.

Om det inte går att bygga egen AI är det nästa logiska steget att införskaffa AI från andra organisationer i ens eget land eller från andra länder. Exempelvis har flera länders försvarsmakter införskaffat AI-modeller [41]. Bland annat har amerikanska militären ingått i ett kontrakt med Open AI [42]. Därtill har Ryssland börjat använda AI för militärt bruk. För Ryssland gäller detta främst AI-modeller som är lättare att få tag på och som sedan kan köras lokalt.

¹ Graphics Processing Unit (GPU) är en processor ursprungligen utvecklad för att hantera grafik. Varje GPU i Colossus har en kapacitet likvärdigt en Nvidia GPU H100.

Det rör sig om AI-modeller med öppna modellvikter såsom Mistral, Qwen och Llama [43].

Lokala AI-modeller med öppna modellvikter kan vara snabbare eftersom de körs lokalt. Men de är typiskt mindre kraftfulla, enligt etablerade prestandajämförelser (eng. *benchmarks*) som exempelvis *Humanities Last Exam* (med resonemang) [43], *Software Engineering Benchmark* (inom mjukvaruutveckling) [44] och *Chatbot Arena* (som utgår från mänsklig preferens) [45].

Externa AI-modeller är alltså kraftfullare men har också begränsningar. Denna typ av modell tillhandahålls genom chatttrutor eller API:er för mjukvaror. När en användare ställer en fråga till modellen skickas frågan till AI-modellens server där beräkningarna utförs. Detta innebär att användaren förlorar kontroll över sina data, vilket kan vara ett problem [32]. Därtill måste användaren förlita sig på extern support och kan inte veta om tjänsten kommer fortsätta finnas. Det finns också en viktig ekonomisk aspekt i att ha råd att betala för önskad AI-användning. Mer om externa molntjänster för AI beskrivs i avsnitt 4.3.

Oavsett om AI-modellen körs lokalt eller externt kan tillverkaren ha implementerat skyddsräcken vid träning av modellen. Leverantören kan använda sådana skyddsräcken för att begränsa prestanda eller funktion (se avsnitt 4.8). Exempelvis verkar kinesiska Deepseek undvika att svara på vissa frågor som är politiskt känsliga för kinesiska staten [32].

4.2.4 Konsekvenser för cyberförsvar

För att utveckla AI-modeller med samma förmåga som de ledande molnbaserade modellerna krävs svårrekryterad spetskompetens samt betydande beräkningsresurser och energiförbrukning. Det är knappast möjligt för cyberförsvaret att själva bygga datacenter i syfte att träna AI-modeller som konkurrerar mot de som byggs av världens största aktörer. För att nå en kostnadseffektiv förmåga bör Försvarsmakten samarbeta med andra.

Det enklaste sättet att få tillgång till kraftfulla AI-modeller är att köpa tillgången av andra länder, vilket inom Nato främst innebär från USA. AI-modellerna kan införskaffas och köras lokalt eller via externa moln. Att köra AI-modeller lokalt betyder normalt lägre krav på egen hårdvara och lägre strömförbrukning men också sämre prestanda. Användning av externa AI-modeller medför bättre prestanda men också stora risker avseende konfidentialitet och tillgänglighet. Dessa aspekter innebär svåra avvägningar för Försvarsmakten.

En gemensam aspekt mellan lokala och externa modeller är att modellernas resultat kan vara felaktiga. Exempelvis kan leverantören ha implementerat någon form av skyddsräcken för att begränsa förmågan eller för att vilseleda användaren. Därtill är det svårt att veta vad AI-modellen egentligen lärt sig och kan.

4.2.5 Rekommendationer till Försvarsmakten

I följande punktlista beskrivs konkreta rekommendationer till Försvarsmakten avseende tillgång till AI-modeller och -beräkningskapacitet.

- R3. Analysera för vad det är lämpligt att använda molnbaserade AI-modeller kontra lokala modeller, exempelvis vilken verksamhet som har kortvariga krav på konfidentialitet.
- R4. Träna på att upptäcka och känna igen skyddsräcken i AI-modeller och där de ger användarna vilseledande svar.

4.3 Molntjänsters utmaningar

Trenden molntjänsters utmaningar handlar om att den säkerhetspolitiska utvecklingen i omvärlden och relaterade striktare regelverk för säkerhet lett till utmaningar som rör molntjänster. Dessa utmaningar rör att organisationer vill ha bättre kontroll över data och tjänster. Den ökade kontrollen kan delvis hanteras med tekniska lösningar som kryptering, delvis genom att skapa redundans med andra molntjänster eller helt migrera till suveräna eller lokala molnlösningar.

4.3.1 Vad säger analysföretagen?

Forrester påstår att den ökande användningen av AI även kommer att öka användningen av moln för AI (*neoclouds*). Forrester menar också att detta delvis kommer ske på bekostnad av de traditionella molnjättarna (*hyperscalers*).

På grund av striktare regelverk för säkerhet och det övergripande införandet av AI lyfter Gartner konfidentiell databehandling (eng. *confidential computing*) som ett viktigt koncept. Konfidentiell databehandling beskrivs som att data skyddas under bearbetning när den sker på en annan aktörs hårdvara. Gartner menar att konfidentiell databehandling kan användas för att nyttja molntjänster utan att äventyra konfidentialiteten. Gartner menar dock också att det i vissa fall är viktigt att flytta hem data till lokala moln. Gartner hävdar att organisationer börjat migrera från de stora globala molntjänsterna till suveräna eller lokala molnlösningar. Denna flytt har Gartner namngett *geopatriation* (eng. *geopatriation*).

Juniper Research pekar på förra årets driftstörningar av molntjänster, såsom Amazon AWS driftstörning i oktober (2025), för att förutse att molntjänstleverantörerna kommer börja bli mer interoperabla med varandra. Detta påstår Juniper Research kommer förenkla för användare av molntjänster att kombinera flera molntjänster i syfte att öka redundansen.

Capgemini menar att det med den ökade användningen av AI kommer ske en övergång till hybrida- och multimolnlösningar.

4.3.2 FOI-forskning

FOI har utfört en del forskning avseende moln. Detta sammanfattas i tabell 12. Dessa rapporter och memon täcker till stor del de trender som analysföretagen lyfter. Geopatriering nämns visserligen inte specifikt i FOI-forskningen, men det gör de aspekter som trenden bygger på i form av rådighetsproblem och andra geopolitiska risker med molnanvändning. Neocloud lyfts däremot inte av någon av dessa rapporter och memon.

Tabell 12. FOI-forskning som tar upp de koncept för molntjänsters utmaningar som lyfts av analysföretagens trendrapporter.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Molnet – möjligheter och begränsningar	FOI-R--3381--SE [46]	2012	Rapporten beskriver möjligheter och risker med molnet, från när molntekniken var i sin linda. Rapporten lyfter bland annat de underliggande risker kring lagar och beroenden till molntjänstleverantören som analysföretagen menar driver geopatriering.
Strukturer för samhällsviktig data	Memo 8677 [47]	2024	Memot lyfter behovet av att behålla nationell kontroll över samhällsviktiga data för att undvika risker med internationella molntjänster som lyder under utländsk jurisdiktion.
En genomgång av säkerhetsincidenter och åtgärder avseende molntjänster	FOI-R--5825--SE [8]	2026	Rapporten analyserar orsaker till säkerhetsincidenter och möjliga åtgärder. Rapporten lyfter även problematik kring molnsuveränitet och rådighet. I rapporten nämns även några tekniker kring konfidentiell databehandling. Rapporten beskriver även olika molndistributionsmodeller såsom multimoln.

4.3.3 Explorativ analys

Moln påverkas av utvecklingen och införandet av AI. Ett exempel är införandet av neocloud som är molntjänstleverantörer som specialiserar sig på att erbjuda GPU som en tjänst (eng. *GPU-as-a-service*). Neoclouds konkurrerar med de större traditionella molntjänstleverantörerna, bland annat genom att vara billigare inom sin specialitet, vilket förmodligen beror på bättre tillgång till GPU:er [48].

Nyttjandet av externa molntjänster medför risker kring konfidentialitet. För att bemöta dessa risker forskas det om konfidentiell databehandling. Konfidentiell databehandling finns redan idag för både CPU:er och GPU:er, såsom AMD SEV-SNP [49], Intel TGX [50], ARM CCA [51] och NVIDIA Hopper architecture [52]. Dock har flera sårbarheter påvisats för dessa tekniker på senare tid, såsom Battering RAM [53], Cache Warp [54], Wire Tap [55], SEV-Step [56] och TEE.fail [57]. Detta gör att det kan ifrågasättas hur väl skyddsteknikerna fungerar mot en kvalificerad hotaktör.

Det har blivit vanligare att oroas för digital suveränitet, vilket också framgår i avsnittet om tillgång till AI-modeller (se avsnitt 4.2). EU publicerade i slutet av 2025 *Cloud Sovereignty Framework*, vilket beskriver suveränitetsmål för att hjälpa organisationer att bedöma till vilken grad en molntjänst är suverän [58]. Suveränitetsmålen är uppdelade i kategorierna strategisk, juridisk, data och AI, operativ, leveranskedjerelaterad, teknisk, ekologisk hållbarhet samt säkerhet och efterlevnad. De nationella cybersäkerhetsmyndigheterna från Frankrike (ANSSI) och Tyskland (BSI) har åtagit sig att skapa en uppsättning av kriterier baserat på ramverket. I åtagandet ingår även att skapa en bedömningsmetod för i vilken utsträckning kriterierna uppfylls [59].

De stora molntjänstleverantörerna säger sig försöka tillhandahålla molntjänster som kan anses suveräna. Några exempel på försök till detta är Amazons *AWS European Sovereign Cloud* [60] samt Googles *Cloud Dedicated* och *Cloud Air-Gapped* (som inte kräver internet) [61]. Noterbart har både Nato och brittiska försvarsdepartementet införskaffat *Google Cloud Air-Gapped* [62] [63]. Suveränitet har också sedan tidigare varit ett fokus för Ryssland som försökt frikoppla sin del av internet från resten av världen, med begränsad framgång [64]. Det finns också forskning som undersöker suveränitet som skapas genom ruttmedvetna (eng. path-aware) nätverksprotokoll. Sådan forskning bedrivs i schweiziska projektet Scion [65] [66].

Under 2025 drabbades flera väletablerade molntjänstleverantörer av omfattande driftstörningar, exempelvis hos Google Cloud Platform i juni [67], Amazon AWS i oktober [68] och Microsoft Azure i december [69]. Framförallt AWS driftstörning i oktober skapade mycket uppmärksamhet då flera välkända och samhällsviktiga tjänster gick ned i flera timmar. För att undvika att tjänster går ned vid driftstörningar är en möjlighet att skapa redundans i form av multimolnlösningar. Dessa kan dock ge ökad komplexitet och attackyta. Enligt en FOI-rapport av Ziegenbein m.fl. är ett stort problem med molntjänster den konfigurationskomplexitet som medföljer och att små felkonfigurationer ofta leder till sårbarheter [8]. Detta förstoras av multimoln, där bland annat kommunikationslänkar mellan molntjänstleverantörerna innebär en ny attackyta. Därtill ökar riskerna för konfidentialiteten om samma data finns hos flera leverantörer.

4.3.4 Konsekvenser för cyberförsvar

Försvarsmakten behöver känna till och hantera de risker med konfidentialitet som följer av att använda externa moln. Tekniskt finns det inga perfekta lösningar. Forskning om konfidentiell databehandling och homomorfisk kryptering kan möjligen på sikt ge rimliga nivåer av säkerhet för data. Det finns också en vanlig avvägning mellan att sträva efter högre konfidentialitet och efter att nå högre tillgänglighet genom att sprida informationen till redundanta platser.

Enligt *Försvarsmaktens hybrida multimolnstrategi* [10] avser Försvarsmakten att bygga multimolnlösningar. Sådana lösningar kan minska problemen som uppstår om ett moln stängs ned av leverantören eller på grund av vissa driftstörningar. Vid införande av multimoln behöver de olika leverantörernas beroenden dock analyseras. Exempelvis är det troligt att alla molnleverantörer i, eller från, samma land möter samma påtryckningar från sin regering. I sin strategi beskriver Försvarsmakten tre typer av suveränitet. Den första är datasuveränitet och handlar bland annat om att kunna välja i vilken jurisdiktion data ska lagras. Den andra är teknologisk suveränitet som bland annat handlar om att kunna geopolitiera. Denna typ av suveränitet beskriver Försvarsmakten som svårast att uppnå idag. Den tredje typen av suveränitet är operativ och handlar bland annat om att ha egen driftpersonal.

4.3.5 Rekommendationer till Försvarsmakten

I följande punktlista beskrivs konkreta rekommendationer till Försvarsmakten avseende molntjänsters utmaningar.

- R5. Analysera för vilka användningsområden som externa moln kan användas genom konfidentiell databehandling eller annan teknik.
- R6. Följ utvecklingen av EU:s Cloud Sovereignty Framework, Natos införskaffade Google Cloud Air-Gapped samt liknande initiativ.

4.4 AI-baserad mjukvaruutveckling

Trenden AI-baserad mjukvaruutveckling handlar om att organisationer ser produktivitetsfördelar och lägre kostnader med att utveckla mjukvara med hjälp av AI. Detta kan innebära att mjukvara kan utvecklas i mindre organisationsdelar och inte behöver köpas in utifrån. Framöver kan mänskliga mjukvaruutvecklare bli effektivare med AI som närmast helautomatiskt skapar mjukvara. När automatiken är stor kallas detta också *vibe coding*.

En risk är att plattformarna inför sårbarheter eller utför något som är skadligt. Därtill kan det finnas begränsningar i hur väl AI kan ta över mer abstrakta uppgifter med att ta fram specifikationer för mjukvaran.

4.4.1 Vad säger analysföretagen?

Gartner tror att AI-baserad mjukvaruutveckling kommer fortsätta växa och ta en större roll i utvecklingsarbetet. De förutser att fram till år 2030 kommer 80 % av utvecklارorganisationerna minska antalet människor som ingår i en utvecklارgrupp och istället tillföra AI. Därtill tror Gartner att allt fler applikationer kommer vara byggda med AI-baserad mjukvaruutveckling.

Capgemini lyfter också att användningen av AI-baserad mjukvaruutveckling kommer fortsätta växa och förutspår att det kommer ändra utvecklارrollen. De tror att rollen kommer att handla mer om det övergripande systemet, arkitekturdelar och att styra AI:n snarare än om manuell testning, front-end-utveckling och paketkonfigurering. Capgemini spår att denna förändring kommer börja 2026, men ger inte konkreta siffror för hur stor förändringen kommer bli eller när övergången kommer vara realiserad. Forrester tror att AI-baserad mjukvaruutveckling kommer att göra seniora mjukvaruutvecklare mer effektiva.

4.4.2 FOI-forskning

FOI har inte utfört riktad forskning om AI-baserad mjukvaruutveckling. Däremot har flera studier angränsat trenden. Dessa studier lyfts fram i tabell 13.

Tabell 13. FOI-forskning som tar upp de koncept för AI-baserad mjukvaruutveckling som lyfts av analysföretagens trendrapporter

Titel	Nummer	År	Sammanfattning av hur forskningen rör trenden
Varför har mjukvaror sårbarheter?	FOI-R--5550--SE [70]	2023	Rapporten tar upp att det finns risker med generativ AI eftersom den kan skriva osäker mjukvara. Rapporten sammanställer en annan studies forskningsförsök som visade att Chat GPT (version 3.5) behövde mycket specifika instruktioner för att skriva säkra program. Därtill var de genererade programmen i många fall långt under normal standard för mänskliga utvecklare. Det beskrivs också att AI-modeller kan vara naiva och anta att användaren inte kommer att mata in ogiltig input, vilket en hotaktör skulle kunna utnyttja.
Framtida gränssnitt - Statusrapport 2021–2023	FOI-R--5568--SE [71]	2024	Rapporten lyfter att generativ AI kan användas för att generera programmeringskod.

Titel	Nummer	År	Sammanfattning av hur forskningen rör trenden
Stora språkmodeller som ett stödverktyg vid militär taktisk planering	FOI-R--5852--SE [35]	2026	Rapporten beskriver att generativ AI kan användas för att ge stöd vid utveckling och då framförallt när koden upprepar tidigare kod.
Militärteknik 2050	FOI-R--5904--SE [3]	2026	Rapporten beskriver att agenter kan användas för utveckling av mjukvara.

4.4.3 Explorativ analys

De ekonomiska och strategiska fördelarna med att effektivisera utvecklingen av mjukvara driver företag till att utforska AI-baserad mjukvaruutveckling. Flera tillverkare av verktyg för AI-baserad mjukvaruutveckling hävdar att de själva använder dessa plattformar:

- Googles VD påstod i oktober 2024 att 25 % av Googles kod var AI-genererad [72].
- Microsofts VD påstod i april 2025 att 30 % av Microsofts kod var AI-genererad [73].
- Anthropic's VD påstod i januari 2026 att mjukvaruutvecklare kan ersättas inom sex till tolv månader [74].

Att skriva kod är dock bara en del av många delar inom mjukvaruutveckling. Bland annat behöver systemet också krävställas och designas, vilket innebär behov av verksamhetskännedom och kunskap om systemkontexten.

Det är än så länge oklart om AI-baserad mjukvaruutveckling är kostnadseffektiv för tillverkning och användning. Tillverkarna av AI-verktygen går enligt uppgift än så länge med förlust i tillverkningen [75]. Detta ställer förmodligen krav på att betydande framsteg görs inom hårdvara eller i hur AI-modellerna fungerar, eller att priserna höjs. Vad gäller tiden det tar att använda plattformarna vid utveckling så har detta utvärderats av forskare. I början av 2025 utfördes ett experiment som tydde på att användandet av AI-baserad mjukvaruutveckling ökade tiden för utvecklingen med ungefär 20 %. Utvecklarna själva bedömde tvärtom att AI-baserad mjukvaruutveckling minskade tidsåtgången med 20 % [76]. Samma forskningsgrupp publicerade under 2026 en ny studie på samma tema. De konstaterar nu att användningen av AI faktiskt ökar produktiviteten med knappt 20 % [77].

Det finns också en viktig aspekt om huruvida det är säkert att använda AI-baserad mjukvaruutveckling. Ur säkerhetsperspektiv är det av speciellt

intresse huruvida det finns sårbarheter i den kod som AI-genereras. Forskning tyder på att kod som är genererad av AI-baserad mjukvaruutveckling kan vara sårbar och svårläslig [78] [79]. Forskning av Kharma m.fl. tyder på att LLM:er ofta misslyckas med att utnyttja moderna säkerhetsfunktioner i programmeringsspråk. Ett exempel som lyfts är att modellerna valt föråldrade metoder för slumpgenerering, trots att säkrare alternativ finns tillgängliga i nyare versioner [78]. Ett annat exempel är lösenordsgenerering, vilket beskrivs i [80]. LLM:er verkar särskilt dåliga på att föreslå lösenord och lösenorden blir ofta snarlika lösenord som föreslagits till många användare. Detta är logiskt med tanke på att modellerna bygger på att prediktera nästa tecken, vilket på sätt och vis är motsatsen till hur ett bra (slumpat) lösenord bör skapas. Exempelvis har många exempel på AI-genererade svaga lösenord hittats på Github.

Forskning från Haque m.fl. beskriver en underliggande problematik i att AI-modellerna speglar de data som modellerna tränats på. Om den publika koden på internet innehåller osäkra mönster, kommer AI:n att reproducera dessa. Haque m.fl. lyfter även att AI-modeller lider av hallucinationer där de genererar logiskt felaktig kod eller hittar på bibliotek. Detta gör att manuell kodgranskning och automatiserade tester är viktigt för att upptäcka dolda fel och säkerställa att koden faktiskt fungerar i komplexa projekt [79]. Förmodligen kommer människor ha svårt att läsa och justera AI-genererad kod, vilket ställer krav på att AI-modeller justerar varandra. Det är en öppen fråga hur väl det kan fungera. Det bör dock lyftas fram att även kod som är skriven av människor lider av sårbarheter. Det är för tidigt att säga om AI-genererad kod är mer eller mindre troligt sårbar än människoskriven kod.

En del praktiska säkerhetsproblem med AI-baserad mjukvaruutveckling har rapporterats. Ett sådant exempel är när AWS AI-baserade utvecklingsplattform Kiro orsakade en driftstörning av ett system i AWS molninfrastruktur till följd av att Kiro valde att starta om ett system i produktion [81]. AWS hävdar dock att orsaken till driftstörningen snarare var handhavandefel och felkonfiguration av Kiro [82].

4.4.4 Konsekvenser för cyberförsvar

Försvarsmakten behöver analysera i vilken utsträckning de kan använda AI-baserad mjukvaruutveckling för att kostnadseffektivt öka sin egen utvecklingsförmåga. Detta måste vägas mot risken att införa sårbarheter i systemen.

En viktig aspekt är att kunna förstå kod skriven av AI, vare sig syftet är att avlägsna eller utnyttja eventuella sårbarheter. Om mänskliga utvecklare inte längre behövs för att skriva kod kommer det på kort sikt vara enklare att få tag i kompetensen att förstå kod. På längre sikt kommer dock nya personer inte att ha den kodskrivarerfarenhet som krävs.

Det är också möjligt att AI-genererad kod kommer vara tillräckligt annorlunda att det i alla fall inte är möjligt att granska utan avsevärt AI-stöd.

En annan möjlighet är att mängden sårbarheter som upptäcks (med AI) kommer bli så stor att mer arbete börjar läggas på tidigare delar av utvecklingsprocessen, där sårbarheter delvis kan byggas bort med arkitektoniska val. Det är också möjligt att AI mest är lämpat för att snabbt ta fram prototyper.

4.4.5 Rekommendationer till Försvarsmakten

I följande punktlista beskrivs konkreta rekommendationer till Försvarsmakten avseende AI-baserad mjukvaruutveckling.

- R7. Träna på att hitta sårbarheter i kod skriven med AI-baserad mjukvaruutveckling, med tanke på de särdrag som präglar AI-kod.
- R8. Analysera hur kodgranskningskompetens ska upprätthållas på sikt om kodskrivning inte längre prioriteras hos universitet och företag.
- R9. Analysera vilka operativa, stödjande och ledande rollers uppgifter som kan försvinna eller flyttas med anledning av AI samt vilka nya roller som kan tillkomma.

4.5 AI för cyberförsvar

Trenden AI för cyberförsvar handlar om att AI allt mer missbrukas av angripare vilket också ökat intresset av proaktivt AI-försvar. AI kan ge styrkehöjande effekter i både försvar och angrepp. Exempelvis kan både försvarare och angripare öka sin förmåga att hitta sårbarheter.

4.5.1 Vad säger analysföretagen?

Proaktiva lösningar kommer enligt Gartner att stå för hälften av pengarna som läggs på mjukvarusäkerhet 2030. Gartner beskriver att detta främst rör sig om en sorts AI-driven förebyggande cybersäkerhet. Som särskild motivering till AI-driven cybersäkerhet ger Gartner (och Deloitte) att det behövs för att möta AI-stödda angrepp. Gartner menar att AI kan användas för att förutse, avbryta och neutralisera cyberangrepp innan de inträffar. Mer konkret beskriver Gartner att lösningarna består av att automatisera moving-target-defence, använda hotunderrättelser för att förutse hoten, täppa till sårbarheter på förhand samt att på avancerat vis obfuskeras och vilseleda angripare.

Deloitte ser AI som en multiplikativ styrkehöjande effekt (eng. *force multiplier*) i både försvar och angrepp. Som exempel på angrepp nämner de att skapa deepfakes, vilket de ser som en etablerad teknik (och som beskrivs mer i avsnitt 4.6). Deloitte menar också att ett framväxande hot är AI-stödda avancerade persistenta hot (AI-APT:er) som kan lära sig snabbt och anpassa sig till

försvararen. Deloitte nämner även möjligheter till att använda AI för riskbedömning, regelefterlevnad och sårbarhetsskanning.

4.5.2 FOI-forskning

FOI har omfattande – men inte heltäckande – forskning om många av de aspekter som ingår i AI för cyberförsvar (och cyberangrepp). Några exempel ges i tabell 14.

Tabell 14. FOI-forskning som tar upp de koncept för AI för cyberförsvar som lyfts av analysföretagens trendrapporter

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Möjligheter för automation av roller inom cybersäkerhetsområdet	FOI Memo 6737 [83]	2019	Memot studerar vilka faktorer som avgör möjligheterna till att automatisera mänskliga cybersäkerhetsuppgifter med datorer. Rapporten drar slutsatsen att de viktigaste faktorerna är tillgång till statistiskt underlag samt att uppgiften inte främst rör människors kreativa eller sociala förmågor.
Verktyg som döljer skadlig kod - En systematisk granskning	FOI-R--5366--SE [84]	2022	Rapporten beskriver bland annat polymorfism, vilket är skadlig kod som kan variera sig för att undgå signaturigenkännande antiviruskydd.
Aktivt skydd i cyberdomänen - En kunskapsöversikt	FOI-R--5797--SE [85]	2025	Rapporten beskriver på övergripande nivå att omvärldens forskning om automatiska skydd och försvar främst rör detektion medan kommersiella lösningar i större utsträckning också täcker in åtgärder som att isolera eller avlägsna hotaktörer.
Verktyg och teknik för CNO-övningar. Slutrapport	FOI-R--5817--SE [86]	2025	Rapporten beskriver FOI-verktøget Lore, vilket mest använts för övningar inom cybersäkerhet. Verktøget har nu även anpassats för penetrationstester samt för att integreras med AI-modeller.

Fler exempel på närliggande FOI-forskning beskrivs i nästa avsnitt om en explorativ analys i ämnet.

4.5.3 Explorativ analys

Vad gäller automatisering av intrångsdetektion konstaterar en FOI-rapport från 2026 [87] att forskningsfronten fortfarande inte lyckats lösa praktiska hinder med att inhämta data och utföra åtgärder. Därtill konstateras att diagnosticerande (eller triage) fått för lite fokus. Andra forskare har under 2026 visat tidiga resultat på att AI kan underlätta den mödosamma översättningen mellan olika detektionsregler [88].

AI kan användas för att identifiera mjukvarusårbarheter. En FOI-rapport från 2026 [89] beskriver tester av olika verktygs förmåga att hitta sårbarheter, däribland två AI-baserade verktyg. AI-verktygen gav tillfredsställande resultat för mindre testfall men bedömdes inte kunna hantera de större testfallen. En slutsats där är att AI-verktyg kan komplettera traditionella statistiska analysverktyg genom att ta vid när de traditionella verktygen pekat ut var sårbarheter verkar finnas.

Både intrångsdetektion och sårbarhetsidentifiering handlar om att hitta olika former av signaler. Amerikanska Rand menar i en rapport från 2026 [90] att AI:s förmåga att dölja är en viktig motkraft mot AI:s förmåga att hitta. Rand beskriver en sorts vilseledning där militära organisationer kan använda AI för att skapa stridsdimma. Därmed kan stora mängder autonoma lockbeten användas för att bygga vilseledningskampanjer.

En översikt över möjligheter att använda AI för angrepp ges i en äldre FOI-rapport från 2020 [91]. Där beskrivs övergripande vilka attacksteg som kan utföras av AI (enligt dåvarande definition) med en bedömning av mognad. En FOI-rapport från 2025 [86] beskriver FOI-verktyget Lore, vilket hittills använts för åtta cybersäkerhetsövningar, en säkerhetstävling samt för forskning kring cyberförsvarsmekanismer. Arbete pågår med att anpassa Lore för penetrationstester, vilka ställer andra krav än cybersäkerhetsövningar. Därtill kommer Lore att integrera stora språkmodeller för att kunna lösa utmaningar som ställer stora krav på kreativitet, och som därmed är svåra att skapa generella lösningar för. Enligt annan forskning är åtminstone vissa angreppstekniker fortfarande svåra för AI-verktyg. Exempelvis visar en forskningsartikel på bättre resultat för traditionella knäckningsverktyg för att knäcka lösenord än de resultat som AI-modeller får på samma uppgift [92].

Utanför forskningsvärlden har de stora AI-bolagen rapporterat hög förmåga som gett oro för offensiv AI-användning. Ett exempel är AI-verktyget Mythos som det rapporterades om i början av 2026. I en forskningsartikel från 2026

beskrivs Mythos som ovanligt bra på att gå från identifierad sårbarhet till lyckat angrepp [93]. Mer om Mythos beskrivs också i ett FOI-memo som rapportförfattarnas projekt tagit fram [94]. Mythos testades under 2026 och om liknande framsteg fortsätter kommer äldre forskning snart att ha svag relevans.

Det finns redan en detaljerad vägledning från det amerikanska standardiseringsorganet Nist i form av IR 8596 om försvar och angrepp med och mot AI [95].

4.5.4 Konsekvenser för cyberförsvar

Det finns många möjligheter med AI för cyberförsvar. Att hantera stora datamängder och agera snabbt ger stora fördelar, åtminstone när angrepp sker brett och numerären är begränsad. Än så länge är den styrkehöjande effekten hos AI främst i form av identifiering av sårbarheter, snarare än i utnyttjandet av sårbarheterna. Det finns potential att lösa även detta senare steg, även om det återstår att se om det går att göra på ett kostnadseffektivt sätt. Utvecklingen har de senaste åren gått mycket snabbt.

AI-verktygen är inte perfekta, vare sig det gäller legitim användning eller deras motståndskraft mot manipulation. Detta innebär behov av dubbelkontroll att inte misstag skett. Exempelvis tenderar framstående AI-verktyg göra grova fel i situationer där det är svårt för människor att förutse att fel kommer uppstå. Om AI:n har behörighet att stänga av användare eller ta bort data kan fel beslut få stora konsekvenser.

Som försvarare gäller det också att avgöra vilka angreppssignaler som är allvarligast samt vilka åtgärder som är relevanta. Här glider försvarsproblemet fort över från en teknisk uppgift till något som verksamhetsansvariga beslutsfattare behöver lösa.

4.5.5 Rekommendationer till Försvarsmakten

I följande punktlista beskrivs konkreta rekommendationer till Försvarsmakten avseende AI för cyberförsvar.

- R10. Skapa förtrogenhet med Nist IR 8596 om försvar och angrepp med och mot AI.
- R11. Analysera hur rätt nivå av tillit till AI ska uppnås och vilka typer av fel som AI kommer göra i vilka försvarssituationer.
- R12. Analysera vilka försvarsuppgifter som är mest lämpade för AI och vilket behov av specialisering som finns för dessa uppgifter.

4.6 Digitalt ursprung

Trenden digitalt ursprung handlar om att organisationer står inför ökade risker med kodmanipulation och AI-driven desinformation. Detta har inneburit ökat fokus på säkerhetsmekanismer som verifierar ursprunget hos mjukvara och data, vad gäller tredjepartskomponenter och AI-genererat innehåll.

Dagens metoder för att detektera AI-genererat innehåll kan delas upp i antingen aktiv eller passiv detektion. Passiv detektion innebär att innehåll analyseras för att hitta vanliga mönster från AI-modeller. Aktiv detektion innebär att innehållet märks antingen vid eller efter generering. Både passiv och aktiv detektion har brister i hur de fungerar och helt tillförlitliga detektionsmetoder saknas i dagsläget.

4.6.1 Vad säger analysföretagen?

Gartner lyfter begreppet *digital proveniens* som innebär att verifiera ursprung hos mjukvara, data och media. De lyfter trenden på grund av ökade risker för att programkod manipuleras, övergivna projekt med öppen källkod samt deepfake-driven desinformation. Gartner tar också upp att ökad reglering, exempelvis delar av *EU AI Act* kan komma att kräva märkning och ursprungsspårning av AI-genererat innehåll.

4.6.2 FOI-forskning

I tabell 15 redovisas FOI-forskning som behandlar att använda AI för att sprida desinformation eller som på annat sätt rör digitalt ursprung.

Tabell 15. FOI-forskning som tar upp de koncept för digitalt ursprung som lyfts av analysföretagens trendrapporter.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Fake news images and genuine resilience	Memo 6870 [96]	2019	Memot beskriver att aktörer kan skapa och sprida desinformation med hjälp av AI. Därtill lyfter memot att mer forskning krävs kring forensiska metoder för att detektera genererat innehåll.
Säkra leveranskedjor för it-system	FOI-R--4851--SE [97]	2019	Rapporten beskriver aspekter som handlar om leveranskedjor av digitala artefakter.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Biometrisk metod för teknisk bevakning	FOI-R--5110--SE [98]	2021	Rapporten beskriver biometrisk metod, vilka känner igen människor utifrån deras biologiska eller beteendebaserade kännetecken.
Generativ AI 2024 – Hur möter vi hotet från deepfakes	Memo 8854 [99]	2025	Ett memo på en sida som beskriver problem med att AI kan skapa innehåll som kan användas för att sprida desinformation.
AI-genererat innehåll och desinformation på sociala medier – en systematisk forskningsöversikt	Memo 9187 [100]	2026	Memot beskriver att den ökade tillgången till AI med förmåga att skapa innehåll gjort det betydligt enklare att producera desinformation snabbare, billigare och i större omfattning. Däremot verkar AI inte ha förändrat desinformationskampanjernas målsättningar eller struktur.

4.6.3 Explorativ analys

AI kan generera innehåll som är svårt att urskilja från mänskligt genererat material. Detta kan vara text, bilder och ljud. Det går bland annat att AI-generera videor i nästintill realtid, givet tillräcklig beräkningskapacitet. Som belysts av FOI-studien från Voytiv m.fl. [100] används denna nya förmåga för att sprida desinformation på sociala medier. I en rapport av Johansson m.fl. [101] beskrevs detektionsmöjligheterna av deepfakes generellt som goda, men att angripare som aktivt försökte kringgå detektionen löpte en högst begränsad upptäcktsrisk.

AI underlättar också för angripare att skapa kopior av webbsidor för att använda vid nätfiske [102]. FOI har studerat nätfiske i litteraturgenomgångar och experiment. Den forskningen har inte varit inriktad på AI, men de övergripande slutsatserna kan ha bäring på effekten av AI för att skapa nätfiske. Enligt forskningsartikeln av Sommestad och Karlzén [103] är det överlag svårt att förutse när mottagare kommer gå på nätfiske. Det verkar som att mottagarna typiskt gör en allmän rimlighetsbedömning, snarare än detaljstuderar meddelandets innehåll. Att anpassa meddelandet till mottagaren verkar inte särskilt kraftfullt och det är därför möjligt att AI främst gör utvecklandet av

nätfiske enklare, snarare än bättre. Kanske är nätfiske redan så framgångsrikt att det inte finns så stor förbättringspotential.

Som beskrivs i FOI-studien av Voytiv m.fl. [100] finns det flera sätt att detektera AI-genererat innehåll. I ett forskningsförsök från studien kunde mänskliga granskare bara identifiera deepfakes av hög kvalitet i ungefär 30 % av bedömningarna. Förutom med människor finns det tekniker som kan delas upp i passiv och aktiv detektion. Aktiv detektion är mer tillförlitlig än passiv detektion [104]. Båda varianterna beskrivs mer nedan.

De passiva detektionsmetoderna bygger antingen på avancerade statistiska mönster eller på stora språkmodeller som tränats för att detektera AI-genererat innehåll. De passiva detektorerna har visat sig icke-robusta mot ändringar i AI-modellen som genererar texten [104] eller där AI-modellen omformulerar texten [105]. Dessutom pågår en kapplöpning med snabb utveckling av både AI-modellerna för generering av innehåll och AI-modellerna för att detektera sådant innehåll.

Aktiv detektion innebär att innehållet märks vid generering eller efteråt. Dessa märkningar sker mot angriparens vilja när angriparen inte har full kontroll över AI-modellen som genererar innehållet. Vattenmärkning faller under denna kategori och är den vanligaste aktiva detektionsmetoden [104]. Som framgår i ett FOI-memo [100] innebär vattenmärkning att AI-modellerna som genererar innehåll även integrerar mönster i det AI-genererade innehållet på ett sätt som människor inte kan urskilja. Dessa vattenmärken är även svåra att ta bort utan att förstöra eller förvränga innehållet. Det pågår arbete med hur leverantörer av AI-modeller som genererar innehåll ska förhålla sig till aktiva detektionsmetoder. Bland annat arbetar EU-kommissionen med att publicera Code of Practice on Marking and Labelling of AI-generated content [106]. Dokumentet beskriver regler för märkning av AI-genererat innehåll för system som genererar audio, bilder, video eller text. Ett FOI-memo [100] lyfter även ett alternativ där endast verifierat verkligt innehåll ska märkas.

4.6.4 Konsekvenser för cyberförsvar

AI-genererat material kan användas för att sprida desinformation både brett och riktat samt både mot organisationer och länder. Desinformation är redan i dag vanligt förekommande i påverkanskampanjer i både krig (exempelvis mot Ukraina) och fred (inklusive mot Sverige). AI kan göra att fler hotaktörer har möjlighet att utföra större kampanjer. Liknande teknik kan också innebära att kvalificerade motståndare kan använda nätfiske oftare och bättre riktat mot enskilda personer eller organisationer.

Detektionsmetoderna för AI-generat innehåll är inte helt tillförlitliga. Det är exempelvis troligt att en kvalificerad hotaktör har förmåga att undgå vattenmärkning av innehåll, eller använder en annan AI-modell där vattenmärkning saknas, alternativt modifierar en vattenmärkande modell eller dess genererade innehåll. Det är också möjligt att inte all märkning kan avlägsnas och att rester kvarlämnas i spår efter angrepp. Det kan göra angrepp lättare att attribuera och därmed identifiera angriparen.

4.6.5 Rekommendationer till Försvarsmakten

I följande punktlista beskrivs konkreta rekommendationer till Försvarsmakten avseende digitalt ursprung.

- R13. Träna på att motstå och förstå AI-baserad social manipulation.
- R14. Följ utvecklingen av märkning av AI-genererat innehåll, exempelvis hur arbete fortlöper med EU:s Code of Practice on Marking and Labelling of AI-generated content.

4.7 Postkvantkryptering

Trenden postkvantkryptering handlar om att organisationer lagt mer fokus på att planera för byten av kryptografiska algoritmer, bland annat för att standardiseringsorgan förespråkar sådana byten. Dessa byten kan behövas för att motstå de nya angrepp som skulle realiseras med kvantdatorer om sådana blev bättre. Dagens asymmetriska kryptoalgoritmer skulle därmed kunna knäckas. En viktig aspekt är också att det kan ta lång tid att byta alla kryptoalgoritmer.

Kryptoalgoritmer kan delas in i två kategorier: symmetriska krypton, där samma nyckel används för både kryptering och dekryptering, och asymmetriska krypton, där krypterings- och dekrypteringsnycklarna är olika. Asymmetriska krypton används bland annat för autentisering av identiteter via digitala signaturer och certifikat. Asymmetriska krypton används också för att etablera en gemensam symmetrisk nyckel. Detta möjliggör upprättandet av autentiserade kommunikationskanaler där symmetriskt krypto, som är mer beräkningseffektivt men kräver en fördistribuerad nyckel, kan användas för att överföra större mängder data. Användningen av båda symmetriska och asymmetriska krypton är mycket utbredd.

4.7.1 Vad säger analysföretagen?

Både Forrester och Juniper Research tror att omställningen till kvantsäkra asymmetriska krypton kommer att utgöra en genomgående teknisk trend de kommande åren.

Här behövs en förklaring av denna rapports författare. Dagens asymmetriska krypton är mycket vanligt använda och bygger sin säkerhet på den antagna höga beräkningskomplexiteten hos två matematiska problem: det diskreta logaritmproblemet och heltalsfaktorisering. En tillräckligt avancerad kvantdator skulle kunna knäcka dessa traditionella asymmetriska kryptosystem. Kvantdatorn gör detta genom att beräkna den diskreta logaritmen, vilket också motsvarar att faktorisera heltal [107].

Forrester uppskattar att övergången till kvantsäkra krypton kommer kräva fem procent av den totala it-säkerhetsbudgeten (utan att närmare specificera detta). Kostnaden motiveras med att det är många system som ska kartläggas, prioriteras och uppgraderas.

Juniper Research yttrar sig mer specifikt angående hybridimplementationer, vilka vanligtvis består av en kvantsäker standard tillsammans med en traditionell algoritm. Enligt Juniper bör dessa prioriteras, eftersom de har potential att leda till riskminimering, bakåtkompatibilitet och operativ stabilitet.

Forrester tror att det är en tidsrymd på tio år innan kvantdatorer utgör ett faktiskt hot, medan Juniper endast pekar på att standardiseringar av algoritmer kommer att påtvinga en snabb övergång.

4.7.2 FOI-forskning

FOI har inte publicerat forskning specifikt inom postkvantkryptering. Däremot finns några publiceringar som anknyter till konceptet, vilka presenteras i tabell 16. Försvarmakten har också egen forskning om detta, se exempelvis [108].

Tabell 16. FOI-forskning som tar upp de koncept för postkvantkryptering som lyfts av analysföretagens trendrapporter

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Kvantteknologier för försvar och säkerhet - Verksamhet och resultat 2024	FOI Memo 8946 [109]	2025	Memot rör tillämpning av kvantteknik och kvantalgoritmer. Detta är relaterad forskning men rör inte specifikt postkvantkryptering.
Skyddsvärd teknik – Årsrapport 2025	FOI Memo 9153 [110]	2026	Memot nämner övergripande kvantdatorers påverkan på traditionella krypton.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Militärteknik 2050	FOI-R--5904--SE [3]	2026	Rapporten lyfter risken att om en tillräckligt kraftig kvantdator kan byggas så kan dagens publika asymmetriska nycklar knäckas. Rapporten lyfter även att Nist har föreslagit en lösning för kvantsäker kryptering.

4.7.3 Explorativ analys

Under flera år har det amerikanska standardiseringsorganet Nist lett ett standardiseringsarbete med fokus på postkvantkryptografiska algoritmer (PQC-algoritmer). Dessa algoritmer baserar sin säkerhet på beräkningskomplexiteten hos matematiska problem som antas vara svåra även för kvantdatorer. Nist har redan valt flera sådana algoritmer, däribland ML-KEM [111] och ML-DSA [112]. Standardiseringen fortgår och det kommer vara ett stort arbete att byta alla algoritmer som behöver bytas. Motsvarande vägledning i Sverige är framtagen av Nationellt cybersäkerhetscenter (NCSC), lett av Försvarets radioanstalt (FRA) [113]. NCSC:s vägledning är delvis baserad på Nists samt har delvis tagits fram i EU-samarbete.

En viktig aspekt är att de traditionella algoritmerna som används i dag visat sig vara motståndiga mot kryptoanalytisk granskning under decennier (förutom mot teoretiska kvantdatorer). Det är för tidigt att säga om PQC-algoritmerna är lika motståndskraftiga i praktiska miljöer som traditionella asymmetriska krypton varit. Därtill tenderar PQC-algoritmerna att ha en större implementationskomplexitet än de traditionella algoritmerna. Denna komplexitet i exempelvis mjukvara kan också innebära nya sårbarheter.

Innan det är möjligt att helt gå över till PQC-algoritmer är Nists och NCSC:s rekommendationer att använda hybridkonfigurationer som är designade för att vara säkra så länge minst en av de ingående algoritmerna är säker [114] [113]. Denna rekommendation följs av företag som Amazon [115], Google [116] och Microsoft [117]. Webbbläsare som Chrome och Firefox använder just nu hybridalgoritmen X25519MLKEM768 som sin standard för nyckelutbyte [118] [119]. Nist beskriver svårigheter med att byta algoritmer och ger vägledning för att upprätthålla en agil förmåga i detta (kryptoagilitet), där byten kan ske samtidigt som interoperabilitet och andra värden upprätthålls [120].

Det är oklart när den första mogna kvantdatoren kommer att finnas, och om det någonsin kommer finnas en sådan. Det finns uppskattningar av hur stor en

kvantdator måste vara för att knäcka traditionella kryptoalgoritmer, bland annat enligt forskning från Försvarsmakten [108]. IBM siktar på 2000 felrättade kvantbitar till år 2033 [121], vilket är tillräckligt för att angripa de minsta kryptoalgoritmerna [122]. Om så stora kvantdatorer verkligen kan skapas är oklart. Men de enorma riskerna motiverar trots allt skyndsamhet, särskilt som en majoritet av världens företag, privatpersoner och organisationer är beroende av asymmetriska krypton för säker kommunikation. Därtill kan angripare samla in krypterade data nu och sedan invänta kvantdatorernas uppkomst för att då kunna dekryptera data. Detta gör det viktigt att påbörja migrationen i god tid innan den första praktiskt användbara kvantdatorn förväntas existera. En enkel formel för detta ges av Moscasc teorem: tiden data måste skyddas plus tiden för att byta algoritm bör hållas kortare än tiden för kvantdatorns intåg [123]. Nist anger att traditionella asymmetriska kryptosystem bör anses föråldrade från och med år 2030, och inte längre tillåtas efter 2035 [114]. NCSC anger också 2035 som slutår medan redan 2030 gäller för särskilt känslig information [113]. Dessa tidsramar är relativt korta, särskilt med tanke på omfattningen av migrationen, vilken innefattar en enorm mängd kryptobibliotek, dedikerad hårdvara, digitala certifikat, blockkedjor med mera. Detta ska dock ställas mot att det finns mycket information som måste vara skyddad under längre tidsram. NCSC rekommenderar att organisationer som omfattas av cybersäkerhetslagen börjar med att lyfta frågan om postkvantkryptering till ledningen samt inventerar den kryptografi som finns i organisationen [113].

Vad gäller symmetriska krypton påverkas de inte i lika stor utsträckning av kvantdatorer som asymmetriska gör. För symmetriska krypton krävs förmodligen en fördubbling av den nyckellängd som anses säker i förhållande till klassiska datorer, vilket typiskt innebär att kvantresistent symmetrisk nycklarna bör ha en nyckellängd på 256 bitar. Det är även möjligt att 128 bitar räcker under lång tid framöver, med tanke på svårigheten att uppnå parallellisering för den kryptoknäckande algoritmen [124].

4.7.4 Konsekvenser för cyberförsvar

Aktörer med långsiktiga ambitioner samlar troligen redan i dag in krypterade data i syfte att kunna dekryptera dessa när tillgång till en mogen kvantdator uppnås (om alls). Försvarsmakten bör ta höjd för detta och redan nu planera för att gå över till kvantsäkra asymmetriska kryptoalgoritmer. De mest prioriterade systemen är förmodligen de med högst exponering och högst krav på konfidentialitet och då särskilt när konfidentialiteten måste upprätthållas under lång tid.

I stort sett alla viktiga aktörer inom kryptografi och säker datoranvändning står inför en snarlik utmaning. Amerikanska standardiseringsorganet Nist brukar vara en mycket inflytelserik aktör inom kryptografi och cyber, vilket även gäller

här. Även svenska NCSC har en liknande vägledning som dock är lite kortare. Det är viktigt att förhålla sig till dessa vägledningar och diskussionerna kring dem, även om Försvarmakten också behöver egen rådighet. Eftersom de nya algoritmerna (och deras implementationer) är otestade i praktiken föreslår Nist och NCSC en övergångsperiod med hybrider som kombinerar traditionella och nya algoritmer. NCSC rekommenderar också att organisationers arbete med övergången inleds med att inventera den kryptografi som finns i organisationen. Dessa sammanställningar kan också komma att vara av nytta för Försvarmakten vid försvar av organisationerna. En viktig aspekt är också att kryptoalgoritmer kan vara säkra i sig men att deras implementationer och användning kan bli fel i praktiken. Sådana praktiska fel kan drabba både Försvarmakten och andra aktörer.

4.7.5 Rekommendationer till Försvarmakten

I följande punktlista beskrivs konkreta rekommendationer till Försvarmakten avseende postkvantkryptering:

- R15. Analysera Nists och NCSC:s vägledningar och notera svåra moment, särskilt i kombination med erfarenheter från tidigare övergångar mellan kryptosystem.

4.8 Säkerhet för stora språkmodeller

Trenden säkerhet för stora språkmodeller handlar om att nya säkerhetsmekanismer införs som svar på de nya säkerhetsproblem som uppstår när AI-modellerna blir allt mer kraftfulla. Dessa säkerhetsproblem består av möjligheter till missbruk av modellerna och även att ta sig runt de säkerhetsmekanismer som leverantörerna infört.

4.8.1 Vad säger analysföretagen?

Gartner lyfter den trend som de kallar AI-säkerhetsplattformar (eng. *AI security platforms*), men som här kallas säkerhet för stora språkmodeller. De nämner inte uttryckligen stora språkmodeller, men har fokus på närliggande risker kring promptinjektion, dataläckage och AI-agenter som rör sig utanför givna ramar. Gartner själva beskriver säkerhetsplattformar för AI som en samling av kontroller för att säkra både egna AI-applikationer och tredje parts AI-tjänster.

Deloitte nämner också flera risker med att använda AI. Vissa av riskerna är specifikt i modellerna, medan andra risker rör data eller omkringliggande applikationer och infrastruktur. En risk som nämns är läckage av träningsdata och de data AI:n använder. En annan risk som nämns är att AI:n får särskilt

utformade förfrågningar som leder till att AI:n överbelastas. Därtill nämns en risk att AI:n kan få för mycket tilldelade behörigheter vars användning innebär ej avsett agerande. Givet riskerna nämner Deloitte också skydd mot att AI missbrukas. Däribland nämns förklarbarhet, datasanering, skyddsräcken, segmentering och penetrationstester (red-teaming).

4.8.2 FOI-forskning

FOI har bedrivit en del forskning kring säkerhet för stora språkmodeller, huvudsakligen de som beskrivs i tabell 17.

Tabell 17. FOI-forskning som tar upp de koncept för säkerhet för stora språkmodeller som lyfts av analysföretagens trendrapporter

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Attacking and Deceiving Military AI Systems	FOI-R--5396--SE [125]	2023	Rapporten beskriver sårbarheter hos AI-modeller. Bland annat beskrivs extrahering av hemlig information samt modellförgiftning.
Large Language Models in Defence: Challenges and Opportunities	FOI-R--5544--SE [27]	2024	Rapporten syftar till att visa hur stora språkmodeller kan tränas för att anpassas till en svensk försvarsdomän. Däremot beskrivs inte hur skydd eller begränsningar läggs till för modeller.
Säker modellindelning - skydd och sårbarheter hos språk-modeller	FOI Memo 8513 [126]	2024	Memot lyfter problematiken med att dela AI-modeller. Memot beskriver därtill ett försök att träna och skydda modellerna.
Large language models potentiella betydelse för utveckling och användning av biologiska vapen	FOI Memo 8883 [127]	2025	Memot beskriver vilka implikationer AI skulle kunna ha på utveckling, produktion och spridning av biologiska stridsmedel. Memot lyfter även kortfattat att många befintliga modeller har säkerhetsfilter (skyddsräcken) men att det finns metoder för att kringgå dessa.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
DeepSeek: Ändamålsenlighet och tillförlitlighet	FOI-R--5787--SE [32]	2026	Rapporten utvärderar kinesiska språkmodellen Deepseek ur ett tillförlitlighets- och ändamålsperspektiv med fokus på användning inom svensk myndighetsutövning.

4.8.3 Explorativ analys

Det finns redan en stor mängd identifierade sårbarheter i språkmodeller. Dessutom har olika branschorganisationer tagit fram taxonomier för att beskriva och hantera sårbarheterna. Några exempel på sådana taxonomier är:

- Owasps topp-10 för stora språkmodeller 2025 [128].
- Mitres Atlas [129].
- Nists AI 100-2e2025 [130].
- Googles Secure AI Framework [131].

Bland de vanliga sårbarhetskategorierna finns promptinjektion [132] [133]. En variant av detta är indirekt promptinjektion, där injektionen döljs i annan data. Ett exempel är en form av vattenhålsangrepp där angriparen skapar en bedräglig webbplats och väntar på att AI ska besöka den. Vid besöket kan angriparen skicka sina kommandon till AI:n [134]. Därtill lyfter Owasps fram språkmodellernas stokastiska beteende som försvarande när det gäller att hindra promptinjektion [128].

Andra exempel på sårbarhetskategorier är dataläckor av olika slag [135] samt modellförgiftning där träningsdata manipuleras [136]. En FOI-rapport [125] beskriver dessa två kategorier. För dataläckage beskrivs särskilda frågor till en AI som därmed svarar med en delmängd av träningsdata. Detta sätt går ut på att få ut (extrahera) en konkret kopia av en delmängd av träningsdata. Vad gäller modellförgiftning beskrivs att AI-modellens träningsdata kan förgiftas varpå AI:n inte fungerar som tillverkaren tänkt. Detta exemplifieras med en AI som klassificerar fordon, där en angripare förgiftar data för att angriparens militära fordon inte ska upptäckas som militära utan istället felaktigt klassificeras som civila.

För att motverka AI-sårbarheter har flera tekniker föreslagits. En vanlig typ av skyddsmekanism är skyddsräcken [137], vilka används för att identifiera och motverka osäkert beteende. Skyddsräcken kan användas av tillverkaren för att kontrollera vad som tillhandahålls till användaren, eller införas av användaren för att inte tappa kontroll över vad AI:n gör. Exempelvis visar en FOI-rapport

från 2026 att den kinesiska språkmodellfamiljen Deepseek [32] tenderar att ge så kallade plakatsvar på frågor ”rörande storpolitik, säkerhetspolitik eller kommunistpartiet i Kina”. Svaren beskrivs utebli eller vara undvikande.

Skyddsräcken kan vara inbakade i modeller genom träning eller implementeras genom externa kontroller då modeller används. Det finns två typer av externa skyddsräcken. Den första typen är att försöka matcha in- och utdata mot en förbestämd mängd ej tillåtna strängar, i ganska enkla textbaserade regler. Den andra typen är att låta andra språkmodeller inspektera och flagga osäkert beteende. Även dessa kontrollmodeller har dock samma fundamentala problem som den modell som ska stödjas. Forskning pekar också på att flera av dessa skyddstekniker är möjliga att kringgå [138] [139] [140]. FOI-forskare har deltagit i myndighetsöverskridande forskning om kringgåenden [141]. En slutsats från den forskningen är att ett kringgående visserligen kan lyckas men att det innebär hög beräkningskostnad eller att AI:ns övriga förmåga sjunker.

Därtill finns det inte enighet i hur AI-säkerheten ska utvärderas, men flera prestandajämförelser (såsom [142] [143] [144]) har föreslagits. Oberoende forskning är relativt omogen och det är generellt oklart i vilken utsträckning de olika säkerhetslösningarna verkligen fungerar. Förmodligen behöver skyddsräcken användas i kombination med traditionella cybersäkerhetsfunktioner såsom sandlådor, åtkomstkontroll och penetrationstester, vilket rekommenderas av Owasp [128].

Användaren kan alltså behöva hantera en AI som inte fungerar som förväntat. Här är förmodligen användaren i behov av förklaringar av AI:ns beteende. FOI har forskat även om detta, så kallad förklarbarhet, se till exempel en äldre rapport från 2019 [145]. Förklarbarhet är tänkt att förklara utdata och algoritmer så att det kan förstås av en annan maskinprocess, eller rentav en människa. Ett problem är att AI kan sakna historiken och förmågan att i efterhand resonera om varför den gjorde något.

4.8.4 Konsekvenser för cyberförsvar

Än så länge är säkerheten för stora språkmodeller relativt låg. Leverantörer kan därför inte säkerställa att deras AI-modeller fungerar som tänkt. Detta innebär i sin tur att användare (kunder) bör använda AI med försiktighet.

Ett sätt att hantera detta för användare är genom att minska exponeringen i systemet som kör AI samt begränsa AI-modellens behörigheter, vilket dock också begränsar nyttan med språkmodellerna. Omkringliggande säkerhetsmekanismer behövs också, men kan även de visa sig vara otillräckliga. Det finns även redan gott om kunskap och praktisk färdighet vad gäller traditionella sårbarheter inom cyberdomänen och detta kan appliceras även på AI. Det finns också taxonomier för AI-sårbarheter.

En annan aspekt är att användare kanske inte har samma mål som leverantörerna. Därför är det ibland en intressant möjlighet för användare att avsiktligt gå runt AI-modellens skyddsräcken för att kunna studera modellen bättre eller få den att fungera bättre för användaren. Att som utvecklare få AI:n att ha rätt mål och hålla fast vid dem är mycket svårt och forskning pågår.

4.8.5 Rekommendationer till Försvarsmakten

Forskarna föreslår följande rekommendationer angående säkerheten för stora språkmodeller:

- R16. Identifiera hur exponering och behörigheter påverkar lämpliga användningsområden för språkmodeller.
- R17. Skapa förtrogenhet med angrepp och skydd i Mitre Atlas och i Owasp's taxonomi.
- R18. Testa olika sätt att kringgå stora språkmodellers säkerhetsmekanismer.

4.9 Fysisk AI

Trenden fysisk AI handlar om att organisationer försöker överföra vinsterna med digital AI till den fysiska omgivningen. Detta innebär att traditionella robotar, smarta byggnader och liknande fysiska system kan bli mer adaptiva och kraftfulla. Resultatet blir en sorts cyberfysiska system där en AI-modell interagerar med sin omgivning. Därmed kombineras AI-modeller med observation av den fysiska världen, följt av förståelse och agerande i den fysiska världen. AI-modellerna kan också lära sig av den fysiska världen och anpassa sig.

Det är svårt att uppnå fysisk AI eftersom den fysiska världen inte är så nedlöst som en typisk chattbot-miljö. En viktig aspekt är att den fysiska världen ger fler möjligheter till promptinjektion. Därtill betyder kopplingen till den fysiska världen att angrepp kan få potentiellt värre konsekvenser.

4.9.1 Vad säger analysföretagen?

Gartner beskriver drivkraften bakom trenden som en vilja från organisationer att ta AI:s digitala effektivisering vidare till den fysiska verkligheten.

Deloitte menar att starka ekonomiska incitament driver industriell användning inom fysisk AI. Deloitte lyfter även hur detta möjliggörs av en kombination av bland annat stora språkmodeller, robotik och specialiserad AI-hårdvara. Samtidigt lyfter samma källa att detta kan innebära risker.

Juniper Research tror på ett genombrott för människolika robotar inom tre år. Som en relaterad trend nämner de motmedel till allt mer autonoma drönare. Deras huvudfokus är inte på cyberdomänen, men ett av motmedlen är att ta

över drönarna med cyberangrepp. Juniper Research beskriver också öppen källkod för smarta byggnader, där AI automatiserar processer för värme och ljus samt optimerar energiförbrukningen. Öppen källkod beskrivs öka interoperabiliteten mellan system. Juniper Research tror att smarta byggnader kommer bli tre gånger så vanliga över de kommande fem åren. Ett problem som de nämner är dock att AI i sig medför en ökad energiförbrukning.

4.9.2 FOI-forskning

FOI bedriver forskning angående fysisk AI, med exempel i tabell 18.

Tabell 18. FOI-forskning som tar upp de koncept för fysisk AI som lyfts av analysföretagens trendrapporter.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
Samverkande IoT-system och databearbetning	FOI-R--4884--SE [146]	2019	Rapporten beskriver uppkopplade tingestar som skapar sakernas internet. En del av detta är smarta städer och byggnader. Sådana beskrivs bland annat kunna hjälpa Försvarsmakten att få bättre lägesbild vid insatser.
Svenska smarta städer. En explorativ studie om förväntad nytta och potentiella sårbarheter.	FOI-R--5117--SE [147]	2021	Rapporten beskriver att utvecklingen av smarta städer drivs av förhoppningar om kostnadsbesparingar och effektiviseringar samt förväntningar på förbättrade lägesbilder. Därtill beskrivs möjliga sårbarheter, vilka främst är tekniska men också sådana som följer av förändrade arbetssätt.
Teknikutveckling UAS - omvärldsanalys och trender	Memo 8336 [148]	2023	Memot fokuserar på obemannade flygande system. Bland annat beskrivs ryska drönare kallade Orlan-10 som ensamma kan användas för artillerield och signalspaning. I grupper av tre kan drönarna användas för telekrig.

Titel	Referens	År	Sammanfattning av hur forskningen rör trenden
CEMA - ett koncept för cyber- och elektromagnetiska aktiviteter	FOI Memo 9268 [149]	2026	Memot beskriver CEMA, vilket utgör integreringen av cyberdomänens och elektromagnetiska spektrumets militära aktiviteter (cyberkrig och telekrig).

4.9.3 Explorativ analys

Det har redan gjorts tidiga framsteg i fysisk AI. Ett exempel ges av Gemini Robotics [150] där språkmodeller utökats till att ta in bilder och text för att sedan låta AI:n agera i form av fysisk förflyttning. Det finns dock frågetecken kring huruvida stora språkmodeller är en bra arkitektur för fysisk AI [151]. Ett centralt problem är bristen på träningsdata [152]. Både Google och Open AI satsar därför på så kallade världsmodeller (eng. world models) [153] [154], vilka är modeller som kan simulera fysiska miljöer för att generera syntetiska träningsdata. En kompletterande teknik till världsmodeller är digitala tvillingar, det vill säga digitala kopior av fysiska miljöer som kontinuerligt synkroniseras med sina verkliga motsvarigheter. I båda fallen är tanken att möjliggöra simulerad robotträning utan krav på verkliga fysiska experiment.

Sammankopplingen mellan enheter i smarta städer innebär bättre möjligheter till avläsning och styrning på distans. Detta kan ge bättre lägesbild men det finns också en risk att data manipuleras av angripare för att lägesbilden ska bli felaktig [146] [147]. Om data lagras i molnet kan det också innebära att lägesbilden kan läsas av molntjänstleverantören [147]. Samma källa nämner också att styrningen av smarta städer kan angripas av motståndare för att försöka påverka trafiksignaler eller stänga av vattenförsörjningen. Smarta städer medför därmed både nya möjligheter och nya utmaningar.

AI har haft lättare för många utmaningar som människor har svårt för (exempelvis beräkningar), men svårare för fysiska utmaningar som människor har lätt för (exempelvis att se eller att lära sig att gå). Detta kan bero på att beräkningssituationerna är mer avgränsade och mest handlar om att ha rätt indata. Därtill har forskning visat potentiella attacker som injicerar instruktioner till fysiska AI-system via deras fysiska omgivning [155] [156]. En sårbarhet här är att information och instruktioner inte är åtskilda. Exempelvis kan en språkmodellbaserad bil stå vid ett rödljus och köra för att den ser en skylt där det står "kör". Här finns också en parallell till chattbottars svårigheter med att resonera om världen. Chattbottarna tenderar att föreslå att användaren går till biltvätten eftersom det är en kort sträcka utan att förstå att bilen behöver tas

med till biltvätten.² Det finns också liknande begränsningar med textförståelse, däribland grundläggande svårigheter i att räkna bokstäver i ord.

4.9.4 Konsekvenser för cyberförsvar

Fysisk AI har potentialen att avsevärt påverka cyberförsvarets uppgifter. Dagens militära drönare kan vara dödliga även när autonomin är relativt enkel. Framtida drönare, baserade på mer avancerad AI, kan dessutom tänkas kombinera fysiska angrepp med cyberangrepp. Därtill förväntas cyberdomänen öka i omfattning när allt fler tingestar får internetanslutningar och kan både övervakas och styras på distans. Särskilt ökar sensorerna i fysisk AI möjligheten till lägesbild av byggnader och städer.

Än så länge är det svårt att implementera den typ av fysisk AI som med avancerade algoritmer styr fysiska tingestar. Träningsdata är begränsade jämfört med den fysiska världens många variabler. Därtill finns stora sårbarheter som följer av svårigheten att särskilja data från instruktioner när indata kommer från den fysiska världen.

4.9.5 Rekommendationer till Försvarsmakten

I följande punktlista beskrivs konkreta rekommendationer till Försvarsmakten avseende fysisk AI:

- R19. Analysera skillnader mellan försvar av fysisk AI och försvar av traditionella cyberfysiska system.
- R20. Analysera varianter av promptinjektion för fysisk AI.
- R21. Öva på CEMA för AI-drönare.

² Se <https://opper.ai/blog/car-wash-test>. Rapportförfattarna har också återskapat detta med frågan ”Jag ska tvätta min bil. Biltvätten är 100 m bort. Ska jag gå eller åka?” vilket gav svar enligt https://chatgpt.com/s/t_6a1efa5248b8819185734321928e7fe1

5 Diskussion

Detta kapitel diskuterar rapportens rekommendationer till cyberförsvaret. Därefter diskuteras hur rekommendationerna förhåller sig till Försvarmaktens strategidokument. Sedan diskuteras huruvida de identifierade trenderna kommer fortsätta. Kapitlet avslutas med att forskningsmässigt ifrågasätta ifall dessa trender verkligen är de mest relevanta.

5.1 Rekommendationer för att hantera trenderna

Rekommendationerna delas här in i: skydd och försvar mot och med AI, organisation samt suveränitet. Rekommendationerna återkommer från resultatkapitlet och har hänvisningar dit via sina rekommendations-ID.

5.1.1 Skydd och försvar mot och med AI

Det tydligaste gemensamma draget mellan trenderna är på teknisk nivå AI. En stor begränsning med AI är svårigheten att se till att AI:ns beteende inte blir skadligt för olika intressenter: användare, leverantör eller andra. Den tekniska utvecklingen fokuserade först på allmän AI-förmåga men har på senare tid även tagit med säkerhet i form av bland annat skyddsräcken. Skyddsräcken leder typiskt till förmågebortfall. Särskilt gäller detta när olika intressenter har olika intressen. Ett skyddsräcke kan då hindra skada mot en intressent på bekostnad av en annan intressents förmåga. Cyberförsvaret ges därför rekommendationerna i tabell 19.

Tabell 19. Rekommendationer för skyddsräcken och sårbarheter.

Rekommendation	ID
Skapa förtrogenhet med angrepp och skydd i Mitre Atlas och i Owasps taxonomi.	R17
Träna på att upptäcka och känna igen skyddsräcken i AI-modeller och där de ger användarna vilseledande svar.	R4
Analysera varianter av promptinjektion för fysisk AI.	R20
Testa olika sätt att kringgå stora språkmodellers säkerhetsmekanismer.	R18
Träna på att hitta sårbarheter i kod skriven med AI-baserad mjukvaruutveckling, med tanke på de särdrag som präglar AI-kod.	R7

Det är också viktigt att hantera de särskilda angrepp som kan ske med hjälp av AI. Cyberförsvaret ges därför rekommendationerna i tabell 20.

Tabell 20. Rekommendationer för angrepp.

Rekommendation	ID
Skapa förtrogenhet med Nist IR 8596 om försvar och angrepp med och mot AI.	R10
Analysera hur multiagenta system kan koordineras utan att kontrollen av agenterna övergår till motståndare.	R1
Träna på att motstå och förstå AI-baserad social manipulation.	R13
Följ utvecklingen av märkning av AI-genererat innehåll, exempelvis hur arbete fortlöper med EU:s Code of Practice on Marking and Labelling of AI-generated content.	R14
Öva på CEMA för AI-drönare.	R21
Analysera skillnader mellan försvar av fysisk AI och försvar av traditionella cyberfysiska system.	R19

5.1.2 Organisation

Om förväntningarna på AI infrias kommer AI möjliggöra nya arbetssätt. Med begränsad framgång inom utvecklingen av AI kommer AI bara vara ett verktyg för vissa tillämpningar. Med större framgång kan AI vara mer autonomt som en agent och rentav ersätta personal eller möjliggöra större numerär. Den som vill nå framgång inom AI behöver ta höjd för dessa möjligheter. Cyberförsvaret ges därför rekommendationerna i tabell 21.

Tabell 21. Rekommendationer för arbetssätt.

Rekommendation	ID
Analysera hur ett AI-baserat cyberförsvar skulle se ut organisatoriskt när det kompletteras av multiagenta system.	R2
Analysera hur kodgranskningskompetens ska upprätthållas på sikt om kodskrivning inte längre prioriteras hos universitet och företag.	R8
Analysera vilka operativa, stödjande och ledande rollers uppgifter som kan försvinna eller flyttas med anledning av AI samt vilka nya roller som kan tillkomma.	R9
Analysera hur rätt nivå av tillit till AI ska uppnås och vilka typer av fel som AI kommer göra i vilka försvarssituationer.	R11
Analysera vilka försvarsuppgifter som är mest lämpade för AI och vilket behov av specialisering som finns för dessa uppgifter.	R12

AI kommer ha störst möjlighet att ta över det arbete som handlar om att behandla stora mängder data och där det finns historiska data över bra beslut (utdata). Kravet på interaktion med människor är också en viktig faktor. Det återstår också att se vilken typ av kreativitet som mogen AI kan uppvisa. Andra viktiga aspekter är tolerabel grad av exponering samt i vilken grad en människa

behöver vara med i steg för att observera, orientera, döma av och agera. I denna loop sker cykler av sekventiell observation, orientering, beslut och agerande.

Bara en trend saknar tydlig koppling till AI och det är postkvantkryptering. Detta är ett komplicerat område med utmaningar från matematik till förändringshantering. En övergång till nya algoritmer kan ge motstånd mot kvantdatorer men också skapa sårbarheter som kan utnyttjas. Cyberförsvaret ges därför rekommendationerna i tabell 22.

Tabell 22. Rekommendationer för hanteringen av postkvantkryptering.

Rekommendation	ID
Analysera Nists och NCSC:s vägledningar och notera svåra moment, särskilt i kombination med erfarenheter från tidigare övergångar mellan kryptosystem.	R15

Samtidigt är kryptering ett integritetskritiskt område där Försvarsmakten behöver egen rådgivning.

5.1.3 Suveränitet

AI och moln köps i stor utsträckning från externa företag. De stora leverantörerna som är relevanta för cyberförsvaret är i huvudsak amerikanska, med enstaka europeiska undantag. Även de underliggande hårdvarorna är för det mesta utomeuropeiska. Å ena sida är det positivt att samverka med de tjänster som erbjuds av andra Nato-länder. Å andra sidan kan beroendet bli alltför ensidigt. Bland flera av trenderna är därför suveränitet en viktig aspekt. Suveränitet innebär att användaren har mer kontroll och rådgivning över data, teknik och driftpersonal. Intresset för suveränitetsaspekter har ökat till följd av striktare lagar och ökade geopolitiska spänningar. Det är dock svårt att nå suveränitet utan att åtminstone tillfälligt förlora viss förmåga. Amerikanska företag har visserligen blivit alltmer måna om att bilda mer fristående juridiska enheter i EU. Men det är tveksamt hur fristående dessa enheter verkligen kan vara. Därtill är en särskilt viktig faktor att vissa AI-leverantörer inte vill tillhandahålla sina tjänster till militära organisationer. Suveränitetsaspekten är komplex. Cyberförsvaret ges därför rekommendationerna i tabell 23.

Tabell 23. Rekommendationer för suveränitet.

Rekommendation	ID
Analysera för vad det är lämpligt att använda molnbaserade AI-modeller kontra lokala modeller, exempelvis vilken verksamhet som har kortvariga krav på konfidentialitet.	R3
Analysera för vilka användningsområden som externa moln kan användas genom konfidentiell databehandling eller annan teknik.	R5
Identifiera hur exponering och behörigheter påverkar lämpliga användningsområden för språkmodeller.	R16

En viktig aspekt är också vilken typ av suveränitet som behövs i respektive fall: för data, teknik eller driftpersonal. En annan viktig aspekt är att samarbete med andra länder kan användas för att sedan kunna uppnå krav på suveränitet. Cyberförsvaret ges därför rekommendationerna i tabell 24.

Tabell 24. Rekommendationer för samarbete.

Rekommendation	ID
Följ utvecklingen av EU:s Cloud Sovereignty Framework, Natos införskaffade Google Cloud Air-Gapped samt liknande initiativ.	R6

Därtill har Försvarmakten redan inlett samarbeten med företag för att bidra till teknikutvecklingen. På det stora hela är dock den svenska försvarssektorn en relativt liten aktör som inte kan påverka när olika tekniker blir användbara.

5.2 Förhållandet till Försvarmaktens strategier

Försvarmakten har sedan tidigare strategier för en del av de trender som lyfts av analysföretagen. Exempelvis lyfter *Försvarmaktens hybrida multimolnstrategi* [10] hur riskerna med digital suveränitet kan minskas. Försvarmakten har även en strategi för AI genom *Försvarmaktens strategi för artificiell intelligens (AI)* [9]. Strategin rör många delar av det som lyfts av analysföretagen. Exempelvis beskriver strategin att AI-system ska utformas och införas med höga krav på säkerhet och resiliens. Även denna strategi lyfter suveränitetsaspekter och då behovet av att skapa en tydlig operativ modell som reglerar hur AI utvecklas, anpassas och implementeras i myndigheten. Strategin beskriver även behov av organisationsanpassningar i form av roll- och befattningsändringar, samarbeten och kompetensspridning. En fundamental ändring i hur man anskaffar och utvecklar system lyfts också för att effektivisera denna process. Strategin är ny vilket också gäller AI i Försvarmakten. Båda strategierna kommer troligen behöva vidareutvecklas med tiden. Därtill behöver strategierna realiseras. Att gå från strategier till verklighet är svårt och medför ett stort arbete. Förhoppningsvis kan realiseringarna stödjas av denna rapports rekommendationer.

5.3 Kommer trenderna fortsätta?

Användningen av AI har redan nått marknaden brett och många företag använder idag AI. Många av trenderna är dessutom så nya att de knappt fått möjlighet att blomma ut. Med tanke på den mängd resurser som läggs på utveckling och framtagning av AI är det troligt att framsteg inom tekniken fortsätter.

Det finns också sådant som talar mot AI som fortsatt trend. De stora kostnaderna för att utveckla ny AI är riskfyllda för företagen. Under 2026 har prissättning för inköp och användning ifrågasatts och leverantörer har börjat ta mer betalt för användningen. Det finns också frågetecken kring huruvida mängden chip, datacenter och strömförsörjning kan växa i den takt AI-utvecklingen kräver. Det är också möjligt att de gängse AI-modellerna är nära att nå sina kognitiva tak och att problem som hallucinationer (där AI hittar på) gör att någon generell eller superb intelligens inte är möjlig med denna teknik. Grundläggande problem finns också med AI:s förmåga att förstå världen, minnas vad den gjort och kunna förklara sitt agerande.

Den enda trenden som inte handlar om AI är postkvantkryptering. Denna trend utgår från att praktiskt användbara kvantdatorer relativt snart förverkligas. Ett sådant förverkligande skulle leda till att en stor del av dagens kryptosystem blev osäkra. Kvantdatorer har funnits i mycket begränsad skala, men det är möjligt att tekniken börjar mogna. Med detta i åtanke har arbete inletts för att redan på förhand byta ut kryptosystemen. Även detta kan bli delvis bortkastade pengar, om kvantdatorerna aldrig blir praktiskt användbara. Det är också möjligt att de nya kryptosystemen tas fram i alltför stor hast och blir mer sårbara för knäckning av traditionella datorer än vad dagens kryptosystem är. Å andra sidan innebär förmodligen övergången till nya algoritmer också att systemen designas om för att underlätta även framtida algoritmbyten. Detta är en sorts framtidssäkring mot ytterligare kryptografiska framsteg.

5.4 Är dessa trender verkligen de mest relevanta?

Trenderna i denna rapport är utvalda för att de lyfts fram av några inflytelserika analysföretag. Analysföretagen har relativt hög samsyn över de trender som tas upp, vilket ökar tilltron till trendernas relevans och påverkan på marknaden. Det kan dock noteras att analysföretagen i vissa av sina trender hänvisar till data från andra analysföretag.³ En brist är att företagens beskrivningar av hur de kommer fram till trenderna är knapphändiga eller helt saknas. Kompletterande trender hade förmodligen kunnat identifierats om andra analysföretag hade använts för identifieringen. Även cybersäkerhetsbolag skulle kunna visa på andra trender. Exempelvis har Fortinet Fortiguard [157] och Google Mandiant [158] mer fokus på konkreta angreppstekniker. Även de kompletterande källorna i de explorativa analyserna av trenderna kan präglas av rapportförfattarnas urval. Detta gäller särskilt som antalet källor och mängden analys är relativt begränsat. Trendernas beskrivningar kan likställas med förstudier där en viktig aspekt varit snabb förmedling till läsare. Sådana upplägg kan vara särskilt fördelaktiga för trendanalys där trenderna snabbt utvecklas eller slutar vara trender. Mer rigorösa analyser riskerar att inte hinna publiceras förrän resultaten är för gamla. En parallell kan dras till att AI-forskare är benägna att publicera sina forskningsartiklar innan de hunnit granskas, med risk för lägre kvalitet på det som läses och citeras av andra.

Ett exempel på svårigheten med förutsägelser är cybersäkerhetsmarknadens reaktion till Anthropic's publikationer om modellen Claude Mythos. Plötsligt målades en bild upp av att AI kunde hitta sårbarheter snabbt och lätt i alla system. Denna bild nyanseras i ett memo som projektet för denna rapport publicerat tidigare [94]. Det är också intressant att enbart Deloitte någorlunda förutsåg utvecklingen, bland de analysföretag vars trendrapporter denna rapport baseras på.

Rapporten kan inte fullt ut svara på om dessa trender är de mest relevanta. Å ena sidan präglas trenderna av AI vilket stämmer väl med generell omvärldsutveckling. Å andra sidan har inte källorna och analyserna i rapporten varit så rigorösa att ett slutligt svar kan ges. Rapporten bör därför ses som en explorativ översikt över omtalade trender, tillsammans med konsekvensanalyser och rekommendationer som sätter trenderna i sammanhanget svenskt cyberförsvar.

³ Se exempelvis trenden "The rise of intelligent ops" i Capgemini's rapport [18] som hänvisar till siffror från Gartner samt trenden "The agentic reality check: Preparing for a silicon-based workforce" i Deloitte's rapport [19] som även den hänvisar till siffror från Gartner.

6 Slutsatser

Rapporten besvarar tre frågor. Nedan besvaras dessa utifrån den övriga rapporten.

6.1 Vilka trender beskrivna av analysföretag är relevanta för cyberförsvaret?

Rapporten har identifierat nio övergripande trender som relevanta för cyberförsvaret. De flesta av trenderna rör AI samt, i mindre mån, moln och postkvantkryptering. De nio trenderna beskrivs kortfattat i tabell 25.

Tabell 25. Trenderna och deras förklaringar.

Övergripande trend	Förklaring
Multiagenta system	Enskilda specialiserade AI-agenter har börjat koordineras för att tillsammans utföra arbetsflöden med flera steg och nå mer komplexa mål.
Tillgång till AI-modeller och AI-beräkningskapacitet	Allt mer AI införs vilket kraftigt ökar behovet av beräkningskapacitet. Organisationer får svårt att få lämplig tillgång till AI-modeller.
Molntjänsters utmaningar	Den säkerhetspolitiska utvecklingen ökar behoven av bättre kontroll över data och tjänster, vilket ställer krav på molntjänsterna.
AI-baserad mjukvaruutveckling	Organisationer ser produktivitetsfördelar och lägre kostnader med att utveckla mjukvara med hjälp av AI.
AI för cyberförsvaret	AI missbrukas allt mer av angripare, vilket också ökar intresset av proaktivt AI-försvaret.
Digitalt ursprung	Organisationer står inför ökade risker med kodmanipulation och AI-driven desinformation, vilket ger behov av bättre verifiering hos tredjepartsmjukvara och AI-genererat innehåll.
Postkvantkryptering	Organisationer har lagt mer fokus på att planera för byten av kryptoalgoritmer, för att motstå de nya angrepp som skulle realiseras med kvantdatorer om sådana blev bättre.
Säkerhet för stora språkmodeller	Nya säkerhetsmekanismer införs som svar på de nya säkerhetsproblem som uppstår när AI-modellerna blir allt mer kraftfulla.
Fysisk AI	Organisationer försöker överföra vinsterna med digital AI till den fysiska omgivningen.

6.2 Vad rekommenderas cyberförsvaret göra för att hantera trenderna?

Trenderna resulterar i möjligheter och hot för cyberförsvaret. För att hantera dessa konsekvenser ger rapporten ett antal rekommendationer till cyberförsvaret. Rekommendationerna indelas i de tre gemensamma kategorierna skydd och försvar mot och med AI, organisatoriska aspekter samt suveränitet. Se tabell 26 för en sammanställning, där ID-numren hänvisar till var i rapporten rekommendationen först gavs.

Tabell 26. Rapportens samlade rekommendationer.

Rekommendation	ID	Kategori
Skapa förtrogenhet med angrepp och skydd i Mitre Atlas och i Owasp's taxonomi.	R17	Skydd och försvar mot och med AI
Träna på att upptäcka och känna igen skyddsräcken i AI-modeller och där de ger användarna vilseledande svar.	R4	
Analysera varianter av promptinjektion för fysisk AI.	R20	
Testa olika sätt att kringgå stora språkmodellers säkerhetsmekanismer.	R18	
Träna på att hitta sårbarheter i kod skriven med AI-baserad mjukvaruutveckling, med tanke på de särdrag som präglar AI-kod.	R7	
Skapa förtrogenhet med Nist IR 8596 om försvar och angrepp med och mot AI.	R10	
Analysera hur multiagenta system kan koordineras utan att kontrollen av agenterna övergår till motståndare.	R1	
Träna på att motstå och förstå AI-baserad social manipulation.	R13	
Följ utvecklingen av märkning av AI-genererat innehåll, exempelvis hur arbete fortlöper med EU:s Code of Practice on Marking and Labelling of AI-generated content.	R14	
Öva på CEMA för AI-drönare.	R21	
Analysera skillnader mellan försvar av fysisk AI och försvar av traditionella cyberfysiska system.	R19	Organisation
Analysera hur ett AI-baserat cyberförsvar skulle se ut organisatoriskt när det kompletteras av multiagenta system.	R2	
Analysera hur kodgranskningskompetens ska upprätthållas på sikt om kodskrivning inte längre prioriteras hos universitet och företag.	R8	
Analysera vilka operativa, stödjande och ledande rollers uppgifter som kan försvinna eller flyttas med anledning av AI samt vilka nya roller som kan tillkomma.	R9	

Rekommendation	ID	Kategori
Analysera hur rätt nivå av tillit till AI ska uppnås och vilka typer av fel som AI kommer göra i vilka försvarssituationer.	R11	
Analysera vilka försvarsuppgifter som är mest lämpade för AI.	R12	
Analysera Nists och NCSC:s vägledningar och notera svåra moment, särskilt i kombination med erfarenheter från tidigare övergångar mellan kryptosystem.	R15	
Analysera för vad det är lämpligt att använda molnbaserade AI-modeller kontra lokala modeller, exempelvis vilken verksamhet som har kortvariga krav på konfidentialitet.	R3	
Analysera för vilka användningsområden som externa moln kan användas genom konfidentiell databehandling eller annan teknik.	R5	Suveränitet
Identifiera hur exponering och behörigheter påverkar lämpliga användningsområden för språkmodeller.	R16	
Följ utvecklingen av EU:s Cloud Sovereignty Framework, Natos införskaffade Google Cloud Air-Gapped samt liknande initiativ.	R6	

6.3 Beskriver trenderna mogen teknik som kommer fortsätta vara relevant?

De ämnen som rapporten behandlat är trender just nu. De flesta av trenderna rör AI, vilket på övergripande nivå redan nått marknaden brett. Användningen varierar dock kraftigt mellan specifika varianter och tillämpningar av AI.

Eftersom enorma summor läggs på att vidareutveckla AI är det rimligt att trenderna med AI också fortsätter de närmaste åren. Satsningarna kan dock också ändras snabbt. Detta beror på att satsningarna hittills kostat mycket och att det finns begränsningar i hur mycket pengar som kan satsas. Det beror också på att de mest populära teknikerna bakom AI har inneboende begränsningar med hallucinationer och svårigheter att skapa en modell av världen. Detta kan göra att generativ AI eventuellt innehar en begränsad potential.

Det finns också en trend som inte rör AI och det är postkvantkryptering. Om större forskningsmässiga framgångar sker med kvantdatorer kommer det finnas behov av postkvantkrypton. Med de hybridinföranden som finns kan sådana krypton delvis ses som mogna. Det är svårt att sia om den exakta utvecklingen av kvantdatorer varför många aktörer behöver ta höjd för kvantdatorerna åtminstone det närmsta decenniet.

7 Referenser

- [1] A. Gudmundson Hunstad, ”Teknisk prognos - Cybersäkerhet”, FOI Memo 4557, 2013.
- [2] E. Dalberg, C. During, G. Kindvall och A. Lindberg, ”Urval av avskanningsämnen 2022 - beskrivning av metod och resultat”, FOI Memo 7868, 2022.
- [3] G. Kindvall m.fl., ”Militärteknik 2050”, FOI-R--5904--SE, 2026.
- [4] I. Källén, ”Militärstrategisk analys i Perspektivstudie 27 - Beskrivning av arbetet och erfarenheter”, FOI-R--5742--SE, 2025.
- [5] M. Hudl Waltin, S. Hellman, H. Lundén och U. Wickenberg Bolin, ”Svaga signaler – Kvantitativ insamling och analys. Verksamhet och resultat 2025”, FOI Memo 9190, 2026.
- [6] A. Maynard, ”What the Rapid Adoption of the "Harness" Metaphor in Artificial Intelligence Reveals About How We Conceptualize Human–AI Relations”, doi:10.2139/ssrn.6352678, 2026.
- [7] E. C. Garrido-Merchán och Á. L. López, ”An Automated Survey of Generative Artificial Intelligence: Large Language Models, Architectures, Protocols, and Applications”, arxiv:2306.02781v4. doi:10.48550/arXiv.2306.02781, 2026.
- [8] J. Ziegenbein och V. Bergström, ”En genomgång av säkerhetsincidenter och åtgärder avseende molntjänster”, FOI-R--5825--SE, 2026.
- [9] Försvarsmakten, ”Försvarsmaktens strategi för artificiell intelligens (AI)”, FM2025-25430:1, 2025.
- [10] Försvarsmakten, ”Försvarsmaktens hybrida multimolnstrategi”, FM2024-3120:26, 2025.
- [11] N. Pollock, R. Williams och L. D'Adderio, ”Figuring out IT markets: How and why industry analysts launch, adjust and abandon categories”, Information and Organization, 2022.
- [12] Juniper Research, ”About Juniper Research”, [Online]. Tillgänglig: <https://www.juniperresearch.com/>. [Använd 21 maj 2026].
- [13] Businesswire, ”Juniper Research Named as Top Three Most Influential Analyst House Globally by Telco Technology Buyers”, 2021. [Online]. Tillgänglig: <https://www.businesswire.com/news/home/20210330005862/en/Juniper-Research->

Named-as-Top-Three-Most-Influential-Analyst-House-Globally-by-Telco-Technology-Buyers. [Använd 21 april 2026].

- [14] A. Murthy och N. P. S., "The Future with New Brand Identity – Success Story of Capgemini: A Case Study", *International Journal of Case Studies in Business IT and Education* Juni 2021. doi:10.47992/IJCSBE.2581.6942.0113, 2021.
- [15] D. Devadiga och S. Aithal, "Deloitte as a Global Professional Services Firm: An Evaluation of Strategy, Innovation, and Market Impact", *Poornaprajna International Journal of Management Education & Social Science* Mars 2026. doi:10.64818/PIJMESS.3107.4626.0044, 2026.
- [16] Gartner, "Top Technology Trends 2026", 2025. [Online]. Tillgänglig: <https://www.gartner.com/en/articles/top-technology-trends-2026>. [Använd 2 mars 2026].
- [17] Juniper Research, "Top 10 Emerging Tech Trends 2026", 2026. [Online]. Tillgänglig: <https://www.juniperresearch.com/marketing/emergingtechtrends2026/>. [Använd 2 mars 2026].
- [18] Capgemini, "Top Tech Trends of 2026", 2026. [Online]. Tillgänglig: https://www.capgemini.com/se-en/wp-content/uploads/sites/20/2026/01/Capgemini_Top_Tech_Trends_Report_2026.pdf. [Använd 2 mars 2026].
- [19] Deloitte, "Tech Trends 2026", 2025. [Online]. Tillgänglig: https://mkto.deloitte.com/rs/712-CNF-326/images/DI_Tech-trends-2026.pdf. [Använd 2 mars 2026].
- [20] Forrester, "2026 Predictions - Technology & Security", 2026. [Online]. Tillgänglig: <https://www.forrester.com/predictions/technology-2026/>. [Använd 2 mars 2026].
- [21] N. Pollock och R. Williams, "Industry analysts – how to conceptualise the distinctive new forms of IT market expertise?", *Accounting, Auditing & Accountability Journal* (2015) 28 (8): 1373–1399. doi:10.1108/AAAJ-03-2015-1989, 2015.
- [22] N. Pollock och R. Williams, "The Business of Expectations: How Promissory Organisations Shape Technology & Innovation", *Social Studies of Science* 40(4):525-548. doi:10.1177/0306312710362275, 2010.
- [23] C. Stokel-Walker, "AI hallucinations appear to be creeping into consulting reports", 14 maj 2026. [Online]. Tillgänglig: <https://sherwood.news/tech/ai-hallucinations-appear-to-be-creeping-into-consulting-reports/>. [Använd 28 maj 2026].
- [24] C. Page, "KPMG's AI report becomes and accidental demo of AI hallucinations", 12 juni 2026. [Online]. Tillgänglig: <https://www.theregister.com/ai-and->

ml/2026/06/12/kpmgs-ai-report-turns-into-a-demo-of-ai-hallucinations/5255029. [Använd 15 juni 2026].

- [25] R. Sapkota, K. I. Roumeliotis och M. Karkee, "Ai agents vs. agentic ai: A conceptual taxonomy, applications and challenges", Information Fusion (103599). doi:10.1016/j.inffus.2025.103599, 2025.
- [26] S. Rydström och R. Johansson, "Utredning av stödverktyg och metodik för utvärdering av AI-metoder", FOI-R--5453--SE, 2023.
- [27] F. Kamrani m.fl., "Large Language Models in Defence: Challenges and Opportunities", FOI-R--5544--SE, 2024.
- [28] J. Schubert, "Förslag på AI teknikutveckling", FOI Memo 8386, 2023.
- [29] S. Finstrella, S. Mariani och F. Zambonelli, "Multi-agent Reinforcement Learning for Cybersecurity: Approaches and Challenges", WOA (pp. 103-118). doi:10.1016/j.iswa.2025.200495, 2024.
- [30] K. Kate m.fl., "LongFuncEval: Measuring the effectiveness of long context models for function calling", arxiv:2505.10570. doi:10.48550/arXiv.2505.10570, 2025.
- [31] Australian Signals Directorate, "Careful adoption of agentic AI services", 2026. [Online]. Tillgänglig: <https://www.cyber.gov.au/business-government/secure-design/artificial-intelligence/careful-adoption-of-agentic-ai-services>. [Använd 21 maj 2026].
- [32] F. Kamrani, L. Kanestad, C. Limér, E. Tjörnhammar, E. Wachtmeister och U. Wickenberg Bolin, "DeepSeek: Ändamålsenlighet och tillförlitlighet", FOI-R--5787--SE, 2026.
- [33] E. Granath och K. Nevall, "Artificiell intelligens och militära operationer", FOI Memo 7807, 2022.
- [34] G. Hillerström, L. Å. Persson, H. Hamrell och H. Ovrén, "Effektiva AI-beräkningar 2024", FOI Memo 8840, 2025.
- [35] Z. Alenljung, V. Lindholm och J. O. Karlsson, "Stora språkmodeller som ett stödverktyg vid militär taktisk planering – Slutsatser från två tillämpade studier av stora språkmodeller", FOI-R--5852--SE, 2026.
- [36] Visual Capitalist, "The world's 20 most powerful AI supercomputers", [Online]. Tillgänglig: <https://www.visualcapitalist.com/the-worlds-most-powerful-ai-supercomputers/>. [Använd 12 mars 2026].

- [37] Xai, "Colossus", 2025. [Online]. Tillgänglig: <https://x.ai/colossus>. [Använd 12 mars 2026].
- [38] S. K. Moore, "AI is a Memory Hog: Its Demand Threatens Whole Categories of Electronics", IEEE Spectrum (vol. 63, no. 4, pp. 28-33, April 2026). doi:10.1109/MSPEC.2026.11476974, 2026.
- [39] Datacenters, "Exploring small modular reactors as the next energy frontier", 2 november 2025. [Online]. Tillgänglig: <https://www.datacenters.com/news/nuclear-powered-data-centers-exploring-small-modular-reactors-as-the-next-energy-frontier>. [Använd 8 maj 2026].
- [40] Amazon Web Services, "Amazon signs agreements for innovative nuclear energy projects to address growing energy demands", 16 oktober 2024. [Online]. Tillgänglig: <https://www.aboutamazon.com/news/sustainability/amazon-nuclear-small-modular-reactor-net-carbon-zero>. [Använd 8 maj 2026].
- [41] U.S Department of War, "War department launches AI acceleration strategy to secure american military AI dominance", 2026. [Online]. Tillgänglig: <https://www.war.gov/News/Releases/Release/Article/4376420/war-department-launches-ai-acceleration-strategy-to-secure-american-military-ai/>. [Använd 13 mars 2026].
- [42] Openai, "Our agreement with the department of war", 2026. [Online]. Tillgänglig: <https://openai.com/sv-SE/index/our-agreement-with-the-department-of-war/>. [Använd 13 mars 2026].
- [43] Center for strategic & international studies, "How Russia is reshaping command and control for AI-enabled warfare", 2026. [Online]. Tillgänglig: <https://www.csis.org/analysis/how-russia-reshaping-command-and-control-ai-enabled-warfare>. [Använd 13 mars 2026].
- [44] C. E. Jimenez m.fl., "Swebench Official leaderboards", 26 februari 2026. [Online]. Tillgänglig: <https://www.swebench.com/>. [Använd 21 maj 2026].
- [45] OproAI, "Chatbot Arena +", 18 maj 2026. [Online]. Tillgänglig: <https://openlm.ai/chatbot-arena/>. [Använd 21 maj 2026].
- [46] H. Karlzén, "Molnet – möjligheter och begränsningar", FOI-R--3381--SE, 2012.
- [47] A. Lioufas och O. Svenonius, "Strukturer för samhällsviktig data", FOI Memo 8677, 2024.
- [48] A. R. d. Almeida, "Idle consumer GPUs as a complement to enterprise hardware for LLM inference: Performance, cost and carbon analysis", 2025 6th International

Artificial Intelligence and Blockchain Conference. doi:10.1145/3775043.3775047, 2026.

- [49] AMD, "AMD Secure Encrypted Virtualization (SEV)", [Online]. Tillgänglig: <https://www.amd.com/en/developer/sev.html>. [Använd 21 maj 2026].
- [50] Intel, "Intel Trust Domain Extensions (Intel TDX)", [Online]. Tillgänglig: <https://www.intel.com/content/www/us/en/developer/tools/trust-domain-extensions/overview.html>. [Använd 21 maj 2026].
- [51] Arm, "Confidential Compute Architecture", [Online]. Tillgänglig: <https://www.arm.com/architecture/security-features/arm-confidential-compute-architecture>. [Använd 21 maj 2026].
- [52] NVIDIA, "NVIDIA Confidential Computing - Securing data and AI models in use", [Online]. Tillgänglig: <https://www.nvidia.com/en-us/data-center/solutions/confidential-computing/>. [Använd 21 maj 2026].
- [53] J. De Meulemeester, D. Oswald, I. Verbauwhede och J. Van Bulck, "Battering RAM: Low-Cost Interposer Attacks on Confidential Computing via Dynamic Memory Aliasing", 47th IEEE Symposium on Security and Privacy, 2026.
- [54] R. Zhang m.fl., "CacheWarp: Software-based Fault Injection using Selective State Reset", 33rd USENIX Security Symposium (USENIX Security 24), 2024.
- [55] A. Seto m.fl., "WireTap: Breaking Server SGX via DRAM Bus Interposition", 2025 SIGSAC Conference on Computer and Communications Security, 2025.
- [56] L. Wilke, J. Wichelmann, A. Rabich och T. Eisenbarth, "SEV-Step: A Single-Stepping Framework for AMD-SEV", IACR Transactions on Cryptographic Hardware and Embedded systems (2024, pp.180-206). doi:10.46586/tches.v2024.i1.180-206, 2023.
- [57] J. Chuang, A. Seto, N. Berrios, S. Van Schaik, C. Garman och D. Genkin, "TEE.fail: Breaking Trusted Execution Environments via DDR5 Memory Bus Interposition", 47th IEEE Symposium on Security and Privacy (2026, pp. 3527-2545). doi:10.1109/SP63933.2026.00101, 2026.
- [58] EU, "Cloud Sovereignty Framework", oktober 2025. [Online]. Tillgänglig: https://commission.europa.eu/document/download/09579818-64a6-4dd5-9577-446ab6219113_en. [Använd 15 december 2025].
- [59] ANSSI och BSI, "Joint Statement by ANSSI and BSI on Cloud Sovereignty Criteria", 17 november 2025. [Online]. Tillgänglig: <https://www.bsi.bund.de/SharedDocs/Downloads/EN/BSI/Publications/ANSSI-BSI-joint-releases/Cloud-Sovereignty-Criteria.pdf>. [Använd 15 december 2025].

- [60] Amazon Web Services, "Europe digital sovereignty", [Online]. Tillgänglig: <https://aws.amazon.com/compliance/europe-digital-sovereignty/>. [Använd 21 maj 2026].
- [61] Google Cloud, "Sovereign Cloud", [Online]. Tillgänglig: <https://cloud.google.com/sovereign-cloud>. [Använd 21 maj 2026].
- [62] Google, "Nato and Google Cloud sign multi-million dollar deal for AI-enabled sovereign cloud", 2025. [Online]. Tillgänglig: <https://www.googlecloudpresscorner.com/2025-11-24-NATO-and-Google-Cloud-Sign-Multi-Million-Dollar-Deal-for-AI-Enabled-Sovereign-Cloud>. [Använd 18 mars 2026].
- [63] Google, "Google Cloud awarded landmark sovereign cloud contract with UK Ministry of Defence", 2025. [Online]. Tillgänglig: <https://www.googlecloudpresscorner.com/2025-09-11-Google-Cloud-Awarded-Landmark-Sovereign-Cloud-Contract-with-UK-Ministry-of-Defence>. [Använd 18 mars 2026].
- [64] C. Vendil Pallin, "Rysslands inte så suveräna internet och kriget i Ukraina", FOI Memo 7925, 2022.
- [65] D. Chagnon, K. Thiry-Atighehchi och G. Chalhoub, "A survey on path validation: Towards digital sovereignty", Computer Networks, 256, 110905. doi:10.1016/j.comnet.2024.110905, 2025.
- [66] B. Trammell, J.-P. Smith och A. Perrig, "Adding path awareness to the Internet architecture", IEEE Internet Computing, 22(2), 96-102. doi:10.1109/MIC.2018.022021673, 2018.
- [67] Google Cloud, "Multiple GCP products are experiencing Service issues", 13 juni 2025. [Online]. Tillgänglig: <https://status.cloud.google.com/incidents/ow5i3PPK96RduMcb1SsW>. [Använd 21 maj 2026].
- [68] Amazon Web Services, "Summary of the Amazon DynamoDB Service Disruption in the Northern Virginia (US-EAST-1) Region", 2025. [Online]. Tillgänglig: <https://aws.amazon.com/message/101925/>. [Använd 21 maj 2026].
- [69] Microsoft Azure, "Post Incident Review (PIR) - Azure Resource Manager - Service management failures affecting Azure Government", december 2025. [Online]. Tillgänglig: https://azure.status.microsoft.com/status/history/?trackingId=ML7_-DWG. [Använd 21 maj 2026].

- [70] H. Karlzén, D. Eidenskog, J. Falkcrona och C. Valassi, ”Varför har mjukvaror sårbarheter?”, FOI-R--5550--SE, 2023.
- [71] B. Levin m.fl., ”Framtida gränssnitt - Statusrapport 2021–2023”, FOI-R--5568--SE, 2024.
- [72] G. McKenna, ”Over 25% of Google’s code is now written by AI—and CEO Sundar Pichai says it’s just the start”, 30 oktober 2024. [Online]. Tillgänglig: <https://fortune.com/2024/10/30/googles-code-ai-sundar-pichai/>. [Använd 21 maj 2026].
- [73] J. Novet, ”Satya Nadella says as much as 30% of Microsoft code is written by AI”, 29 april 2025. [Online]. Tillgänglig: <https://www.cnbc.com/2025/04/29/satya-nadella-says-as-much-as-30percent-of-microsoft-code-is-written-by-ai.html>. [Använd 21 maj 2026].
- [74] J. Small, ”Anthropic CEO Says AI Could Replace Software Engineers in 6 to 12 Months”, 21 januari 2026. [Online]. Tillgänglig: <https://www.entrepreneur.com/business-news/ai-ceo-says-software-engineers-could-be-replaced-in-months/502087>. [Använd 21 maj 2026].
- [75] R. Karma, ”Something Ominous Is Happening in the AI Economy”, 10 december 2025. [Online]. Tillgänglig: <https://www.theatlantic.com/economy/2025/12/nvidia-ai-financing-deals/685197/>. [Använd 11 maj 2026].
- [76] J. Becker, N. Rush, E. Barnes och D. Rein, ”Measuring the impact of early-2025 AI on experienced open-source developer productivity”, arXiv:2507.09089. doi:10.48550/arXiv.2507.09089, 2025.
- [77] J. Becker, N. Rush, T. Cunningham, D. Rein och K. Mahamud, ”We are changing our developer productivity experiment design”, 24 februari 2026. [Online]. Tillgänglig: <https://metr.org/blog/2026-02-24-uplift-update/#selected-developer-quotes>. [Använd 22 maj 2026].
- [78] M. Kharma, S. Choi, M. AlKhanafseh och D. Mohaisen, ”Security and quality in LLM-generated code: a multi-language, multi-model analysis”, IEEE Transactions on Dependable and Secure Computing (vol. 23, no. 3, pp. 7300-7314). doi:10.1109/TDSC.2026.3672745, 2025.
- [79] A. Haque m.fl., ”Exploring hallucinations and security risks in AI-assisted software development with insights for LLM deployment”, 2025 Sixth International Conference on Intelligent Data Science Technologies and Applications (pp. 57-64). doi:10.1109/IDSTA66210.2025.11202778, 2025.

- [80] Irregular, "Vibe Password Generation: Predictable by Design", 18 februari 2026. [Online]. Tillgänglig: <https://www.irregular.com/publications/vibe-password-generation>. [Använd 28 maj 2026].
- [81] Financial Times, "Amazon service was taken down by AI coding bot", 2026. [Online]. Tillgänglig: <https://www.ft.com/content/00c282de-ed14-4acd-a948-bc8d6bdb339d>. [Använd 19 mars 2026].
- [82] Amazon Web Services, "Correcting the financial times report about AWS, Kiro and AI", 2026. [Online]. Tillgänglig: <https://www.aboutamazon.com/news/aws/aws-service-outage-ai-bot-kiro>. [Använd 19 mars 2026].
- [83] T. Sommestad, J. Brynielsson och S. Varga, "Möjligheter för automation av roller inom cybersäkerhetsområdet", FOI Memo 6737, 2019.
- [84] H. Holm, E. Hyllienmark och M. Persson, "Verktyg som döljer skadlig kod - En systematisk granskning", FOI-R--5366--SE, 2022.
- [85] H. Karlzén, J. Ziegenbein och C. Vestlund, "Aktivt skydd i cyberdomänen. En kunskapsöversikt", FOI-R--5797--SE, 2025.
- [86] H. Holm, "Verktyg och teknik för CNO-övningar - Slutrapport", FOI-R--5817--SE, 2025.
- [87] T. Sommestad och H. Karlzén, "Automatisk incidenthantering - Erfarenheter från labbförsök", FOI-R--5878--SE, 2026.
- [88] M. Xu m.fl., "ARuleCon: Agentic Security Rule Conversion", In Proceedings of the ACM Web Conference 2026 (pp. 3123-3134). doi:10.1145/3774904.3792458, 2026.
- [89] A. Rosén och L. Nyholm, "Tester av verktyg som identifierar mjukvarusårbarheter", FOI-R--XXXX--SE, I tryck.
- [90] Z. Burdette m.fl., "How Artificial Intelligence Could Reshape Four Essential Competitions in Future Warfare", 22 januari 2026. [Online]. Tillgänglig: https://www.rand.org/pubs/research_reports/RRA4316-1.html. [Använd 1 juni 2026].
- [91] E. Zouave, M. Bruce, K. Colde, M. Jaitner, I. Rodhe och T. Gustafsson, "Artificially intelligent cyberattacks", FOI-R--4947--SE, 2020.
- [92] M. A. Rehman, S. I. A. Shah, A. Anwar, N. Islam och H. Khan, "(2025). When Intelligence Fails: An Empirical Study on Why LLMs Struggle with Password Cracking. arXiv preprint", arXiv:2510.17884. doi:10.48550/arXiv.2510.17884, 2025.

- [93] Z. Wang m.fl., "ExploitGym: Can AI Agents Turn Security Vulnerabilities into Real Attacks?", arXiv preprint arXiv:2605.11086. doi:10.48550/arXiv.2605.11086, 2026.
- [94] V. Bergström och J. Ziegenbein, "Mythos - Är ryktena om cybersäkerhetens död överdrivna?", FOI Memo 9314, 2026.
- [95] K. Megas m.fl., "Cybersecurity Framework Profile for Artificial Intelligence (Cyber AI Profile)", NIST IR 8596 (Initial preliminary draft). doi:10.6028/NIST.IR.8596.iprd, 2025.
- [96] N. Wadströmer, D. Gustafsson och P. Thunholm, "Fake news images and genuine resilience", FOI Memo 6870, 2019.
- [97] D. Eidenskog, C. Bildsten och B. Endres, "Säkra leveranskedjor för it-system", FOI-R--4851--SE, 2019.
- [98] H. Karlzén, C. Valassi, J. Bengtsson och A. Gudmundson Hunstad, "Biometrisk metod för teknisk bevakning", FOI-R--5110--SE, 2021.
- [99] S. Öberg, "Generativ AI 2024 – Hur möter vi hotet från deepfakes", FOI Memo 8854, 2025.
- [100] S.-Y. Voytiv och M. Svahn, "AI-genererat innehåll och desinformation på sociala medier – en systematisk forskningsöversikt", FOI Memo 9187, 2026.
- [101] F. Johansson m.fl., "Detection of Fabricated Media", FOI-R--5132--SE, 2021.
- [102] M. S. Liaqat, G. Mumtaz, N. Rasheed och Z. Mubeen, "Exploring Phishing Attacks in the AI Age: A Comprehensive Literature Review", Journal of Computing & Biomedical Informatics, 7(02). 2024.
- [103] T. Sommestad och H. Karlzén, "The unpredictability of phishing susceptibility: results from a repeated measures experiment", Journal of Cybersecurity, 10(1), tyae021. doi:10.1093/cybsec/tyae021, 2024.
- [104] X. Lingyun, L. Nian, L. Yuling och H. Jiayong, "AI-Generated Text Detection: A Comprehensive Review of Active and Passive Approaches", Computers, Materials & Continua (Volume 86, Issue 3). doi:10.32604/cmc.2025.073347, 2026.
- [105] Y. Cheng, V. S. Sadasivan, M. Saberi, S. Saha och S. Feizi, "Adversarial Paraphrasing: A Universal Attack for Humanizing AI-Generated Text", arxiv:2506.07001. doi:10.48550/arXiv.2506.07001, 2025.
- [106] Europeiska kommissionen, "Second Draft - Code of practice on transparency of AI-generated content", 5 mars 2026. [Online]. Tillgänglig: <https://digital->

strategy.ec.europa.eu/en/library/commission-publishes-second-draft-code-practice-marking-and-labelling-ai-generated-content. [Använd 11 maj 2026].

- [107] E. Bach, "Discrete Logarithms and Factoring", EECS Department, University of California, Berkeley Technical Report No. UCB/CSD-84-186, 1984.
- [108] C. Gidney och M. Ekerå, "How to factor 2048 bit RSA integers in 8 hours using 20 million noisy qubits", *Quantum*, 5, 433. doi:10.22331/q-2021-04-15-433, 2021.
- [109] P. Jonsson, "Kvantteknologier för försvar och säkerhet - Verksamhet och resultat 2024", FOI Memo 8946, 2025.
- [110] G. Kindvall och D. Sundelin, "Skyddsvärd teknik – Årsrapport 2025", FOI Memo 9153, 2026.
- [111] National Institute of Standards and Technology, "Module-lattice-based key-encapsulation mechanism standard", Federal Information Processing Standards Publication 203. doi:10.6028/NIST.FIPS.203, 2024.
- [112] National Institute of Standards and Technology, "Module-lattice-based digital signature standard (FIPS 204)", Federal Information Processing Standards Publication 204. doi:10.6028/NIST.FIPS.204, 2024.
- [113] NCSC, "Nationella rekommendationer för övergången till kvantsäker kryptografi", 2026.
- [114] National Institute of Standards and Technology, "Transition to Post-Quantum", NIST Internal Report 8547. doi:10.6028/NIST.IR.8547.ipd, 2024.
- [115] Google Cloud Security Team, "Announcing Quantum-Safe Key Encapsulation Mechanisms in Cloud KMS", oktober 2025. [Online]. Tillgänglig: <https://cloud.google.com/blog/products/identity-security/announcing-quantum-safe-key-encapsulation-mechanisms-in-cloud-kms>. [Använd 28 april 2026].
- [116] AWS Cryptography Team, "AWS Post-Quantum Cryptography Migration Plan", 5 december 2024. [Online]. Tillgänglig: <https://aws.amazon.com/blogs/security/aws-post-quantum-cryptography-migration-plan/>. [Använd 28 april 2026].
- [117] Microsoft Security Team, "Post-Quantum Cryptography Comes to Windows Insiders and Linux", maj 2025. [Online]. Tillgänglig: <https://techcommunity.microsoft.com/blog/microsoft-security-blog/post-quantum-cryptography-comes-to-windows-insiders-and-linux/4413803>. [Använd 28 april 2].
- [118] D. Adrian, B. Beck, D. Benjamin och D. O'Brien, "Advancing Our Amazing Bet on Asymmetric Cryptography", 23 maj 2024. [Online]. Tillgänglig:

<https://blog.chromium.org/2024/05/advancing-our-amazing-bet-on-asymmetric.html>. [Använd 28 april 2026].

- [119] J. Schanck, "Let mlkem768x25519 support for desktop TLS connections ride the trains", 16 september 2024. [Online]. Tillgänglig: https://bugzilla.mozilla.org/show_bug.cgi?id=1919097. [Använd 28 april 2026].
- [120] E. Barker m.fl., "Considerations for Achieving Crypto Agility - Strategies and Practices", NIST Cybersecurity White Paper. NISC CSWP 39. doi:10.6028/NIST.CSWP.39, 2025.
- [121] IBM, "The future of computing is quantum-centric", [Online]. Tillgänglig: <https://www.ibm.com/roadmaps/quantum/>. [Använd 8 juni 2026].
- [122] R. Babbush m.fl., "Securing Elliptic Curve Cryptocurrencies against Quantum Vulnerabilities: Resource Estimates and Mitigations", arxiv:2603.28846. doi:10.48550/arXiv.2603.28846, 2026.
- [123] M. Mosca, "Cybersecurity in an Era with Quantum - Will We Be Ready?", IEEE Security & Privacy (Volume: 16, Issue: 5). doi:10.1109/MSP.2018.3761723, 2018.
- [124] NIST, "Post-Quantum Cryptography", 11 december 2025. [Online]. Tillgänglig: <https://csrc.nist.gov/Projects/post-quantum-cryptography/faqs>. [Använd 8 juni 2026].
- [125] F. Kamrani m.fl., "Attacking and Deceiving Military AI Systems", FOI-R--5396--SE, 2023.
- [126] A. Alness Borg, L. Kanestad, S. Lindquist, N. Normelli och B. Pelzer, "Säker modelldelning - skydd och sårbarheter hos språkmodeller", FOI Memo 8513, 2024.
- [127] P. Stenberg och P. Wikström, "Large language models potentiella betydelse för utveckling och användning av biologiska vapen", FOI Memo 8883, 2025.
- [128] OWASP Foundation, "OWASP Top 10 for LLM Applications 2025", 2025. [Online]. Tillgänglig: <https://owasp.org/www-project-top-10-for-large-language-model-applications/assets/PDF/OWASP-Top-10-for-LLMs-v2025.pdf>. [Använd 7 april 2026].
- [129] The MITRE Corporation, "MITRE ATLAS", [Online]. Tillgänglig: <https://atlas.mitre.org/>. [Använd 7 april 2026].
- [130] A. Vassilev, A. Oprea, A. Fordyce, H. Anderson, X. Davies och M. Hamin, "Adversarial Machine Learning: A Taxonomy and Terminology of Attacks and

- Mitigations”, National Institute of Standards and Technology AI 100-2e2025. doi:10.6028/NIST.AI.100-2e2025, 2025.
- [131] Google, ”Google Secure AI Framework”, [Online]. Tillgänglig: <https://saif.google/secure-ai-framework/risks>. [Använd 1 juni 2026].
- [132] A. Zou, Z. Wang, N. Carlini och M. Nasr, ”Universal and Transferable Adversarial Attacks on Aligned Language Models”, arxiv: 2307.15043v2. doi:10.48550/arXiv.2307.15043, 2023.
- [133] P. Wang m.fl., ”The landscape of prompt injection threats in llm agents: From taxonomy to analysis”, arxiv:2602.10453. doi:10.48550/arXiv.2602.10453, 2026.
- [134] Mitre, ”AI ClickFix: Hijacking Computer-Use Agents Using ClickFix”, 31 mars 2026. [Online]. Tillgänglig: <https://atlas.mitre.org/studies/AML.CS0055>. [Använd 28 maj 2026].
- [135] N. Carlini m.fl., ”Extracting Training Data from Large Language Models”, Proceedings of the 30th USENIX Security Symposium (USENIX Security 21). doi:10.48550/arXiv.2012.07805, 2020.
- [136] A. Souly m.fl., ”Poisoning Attacks on LLMs Require a Near-constant Number of Poison Samples”, arXiv:2510.07192. doi:10.48550/arXiv.2510.07192, 2025.
- [137] C. Jarvis, ”How to implement LLM guardrails - OpenAI”, [Online]. Tillgänglig: https://developers.openai.com/cookbook/examples/how_to_use_guardrails. [Använd 7 april 2026].
- [138] M. Nasr m.fl., ”The Attacker Moves Second: Stronger Adaptive Attacks Bypass Defenses Against Llm Jailbreaks and Prompt Injections”, arxiv:2510.09023. doi:10.48550/arXiv.2510.09023, 2025.
- [139] R. J. Young, ”Evaluating the Robustness of Large Language Model Safety Guardrails Against Adversarial Attacks”, 27 November 2025. [Online]. Tillgänglig: <https://doi.org/10.48550/arXiv.2511.22047>. [Använd 7 april 2026].
- [140] W. Hackett, L. Birch, S. Trawicki, N. Suri och P. Garraghan, ”Bypassing Prompt Injection and Jailbreak Detection in LLM Guardrails”, Proceedings of the The First Workshop on LLM Security (2025 aug., pp.101-114). doi:10.48550/arXiv.2504.11168, 2025.
- [141] J. Rosengren, J. Brynielsson, F. Johansson och P. Jonell, ”Jailbreaking Large Language Models: Safety Alignment, Response Quality, Computational Cost”, I 2025 International Conference on Machine Learning and Applications (ICMLA) (pp. 11). doi:10.1109/ICMLA66185.2025.00172, 2025.

- [142] D. Maliugina, "10 LLM safety and bias benchmarks", 17 september 2025. [Online]. Tillgänglig: <https://www.evidentlyai.com/blog/llm-safety-bias-benchmarks>. [Använd 8 april 2025].
- [143] HarmBench, "HarmBench: Standard for LLM Safety Auditing", 21 januari 2026. [Online]. Tillgänglig: <https://www.emergentmind.com/topics/harmbench>. [Använd 8 april 2026].
- [144] S. Tedeschi m.fl., "ALERT: A Comprehensive Benchmark for Assessing Large Language Models' Safety through Red Teaming", arxiv:2404.08676. doi:10.48550/arXiv.2404.08676, 2024.
- [145] L. Luotsinen, D. Oskarsson, P. Svenmarck och U. Wickenberg Bolin, "Explainable Artificial Intelligence: Exploring XAI Techniques in Military Deep Learning Applications", FOI-R--4849--SE, 2019.
- [146] R. Johansson, M. Andersson, P. Hörling och F. Kamrani, "Samverkande IoT-system och databearbetning", FOI-R--4884--SE, 2019.
- [147] J. Bengtsson och L. Mickelsson, "Svenska smarta städer - En explorativ studie om förväntad nytta och potentiella sårbarheter.", FOI-R--5117--SE, 2021.
- [148] M. Norgren och A. Samimi Johansson, "Teknikutveckling UAS - omvärldsanalys och trender", FOI Memo 8336, 2023.
- [149] H. Karlzén, "CEMA – ett koncept för cyber- och elektromagnetiska aktiviteter", FOI Memo 9268, 2026.
- [150] Gemini Robotics Team, "Gemini Robotics: Bringing AI into the Physical World", arxiv: 2503.20020. doi:10.48550/arXiv.2503.20020, 2025.
- [151] M. Assran m.fl., "V-JEPA 2: Self-Supervised Video Models Enable Understanding, Prediction and Planning", arxiv:2506.09985. doi:10.48550/arXiv.2506.09985, 2025.
- [152] Y. Liu m.fl., "Aligning Cyber Space With Physical World: A Comprehensive Survey on Embodied AI", IEEE/ASME Transactions on Mechatronics (vol. 30, no. 6, pp. 7253-7274). doi:10.1109/TMECH.2025.3574943, 2025.
- [153] J. Parker-Holder och S. Fruchter, "Genie 3: A new frontier for world models", 5 augusti 2025. [Online]. Tillgänglig: <https://deepmind.google/blog/genie-3-a-new-frontier-for-world-models/>. [Använd 8 april 2026].
- [154] OpenAI, "Sora 2 is here", 30 september 2025. [Online]. Tillgänglig: <https://openai.com/index/sora-2/>. [Använd 28 april 2026].

- [155] L. Burbano m.fl., "CHAI: Command Hijacking against embodied AI", arxiv:2510.00181. doi:10.48550/arXiv.2510.00181, 2025.
- [156] C. Li m.fl., "The Shawshank Redemption of Embodied AI: Understanding and Benchmarking Indirect Environmental Jailbreaks", arxiv:2511.16347. doi:10.48550/arXiv.2511.16347, 2025.
- [157] Fortinet, "Global Threat Landscape Report", 2026. [Online]. Tillgänglig: <https://www.fortinet.com/resources/reports/threat-landscape-report>. [Använd 1 juni 2026].
- [158] Mandiant, "M-Trends 2026 Report: Real-world investigations and actionable defense insights", 2026. [Online]. Tillgänglig: <https://cloud.google.com/security/resources/m-trends>. [Använd 1 juni 2026].

Bilaga A Extraherade trender

Trendnamn hos analysföretaget	Analysföretag	Inkluderad i övergripande trend
A quarter of CIOs will be asked to bail out business-led AI failures in their organisations	Forrester	<i>(exkluderad)</i>
Enterprises will delay 25% of AI spend into 2027	Forrester	<i>(exkluderad)</i>
Wireless EV Charging: Accelerated infrastructure rollouts drive mass adoption	Juniper Research	<i>(exkluderad)</i>
Cutting through the noise: tech signals worth tracking as AI advances	Deloitte	<i>(exkluderad)</i>
Multiagent systems	Gartner	Multiagenta system
Domain-specific language models	Gartner	Multiagenta system
Multiagent systems: Enterprises invest in domain-specific agents	Juniper Research	Multiagenta system
The agentic reality check: Preparing for a silicon-based workforce	Deloitte	Multiagenta system
The great rebuild: Architecting an AI-native tech organization	Deloitte	Multiagenta system
The year of truth for AI	Capgemini	Multiagenta system
The rise of intelligent ops	Capgemini	Multiagenta system
AI supercomputing platforms	Gartner	Tillgång till AI-modeller och -beräkningskapacitet
Neuromorphic computing: Commercial chipsets that address ai bottlenecks to launch in 2026	Juniper Research	Tillgång till AI-modeller och -beräkningskapacitet
Microfluidics to receive growing interest as next-generation cooling system for AI chips	Juniper Research	Tillgång till AI-modeller och -beräkningskapacitet
Small modular reactors: regulatory approvals open potential disruptive impact on energy generation	Juniper Research	Tillgång till AI-modeller och -beräkningskapacitet

Trendnamn hos analysföretaget	Analysföretag	Inkluderad i övergripande trend
The AI infrastructure reckoning: Optimizing compute strategy in the age of inference economics	Deloitte	Tillgång till AI-modeller och -beräkningskapacitet
The borderless paradox of technological sovereignty	Capgemini	Tillgång till AI-modeller och -beräkningskapacitet
Confidential computing	Gartner	Molntjänsters utmaningar
Geopatriation	Gartner	Molntjänsters utmaningar
Neoclouds will grab \$20 billion in revenue, eroding hyperscaler dominance in genAI	Forrester	Molntjänsters utmaningar
Multi-cloud models: 2025 outages bring focus on more resilience in 2026	Juniper Research	Molntjänsters utmaningar
Cloud 3.0 - all flavors of cloud	Capgemini	Molntjänsters utmaningar
AI-native development platforms	Gartner	AI-baserad mjukvaruutveckling
The time to fill developer positions will double	Forrester	AI-baserad mjukvaruutveckling
AI is eating software	Capgemini	AI-baserad mjukvaruutveckling
Preemptive cybersecurity	Gartner	AI för cyberförsvar
The AI dilemma: securing and leveraging AI for cyber defense	Deloitte	AI för cyberförsvar
Digital provenance	Gartner	Digitalt ursprung
Quantum security spending will exceed 5% of overall IT security budget	Forrester	Postkvantkryptering
Post-quantum cryptography: Standardisation to drive hybrid deployment models	Juniper Research	Postkvantkryptering
AI security platforms	Gartner	Säkerhet för stora språkmodeller
The AI dilemma: securing and leveraging AI for cyber defense	Deloitte	Säkerhet för stora språkmodeller
Physical AI	Gartner	Fysisk AI

Trendnamn hos analysföretaget	Analysföretag	Inkluderad i övergripande trend
Physical AI: Substantial advances in humanoid robotics expected in next 3 years	Juniper Research	Fysisk AI
Counter-drone technology: Growing threats necessitate new technologies	Juniper Research	Fysisk AI
AI goes physical: navigating the convergence of AI and robotics	Deloitte	Fysisk AI
Open-source smart buildings: Interoperable platforms drive market growth amidst growing energy demands	Juniper Research	Fysisk AI

