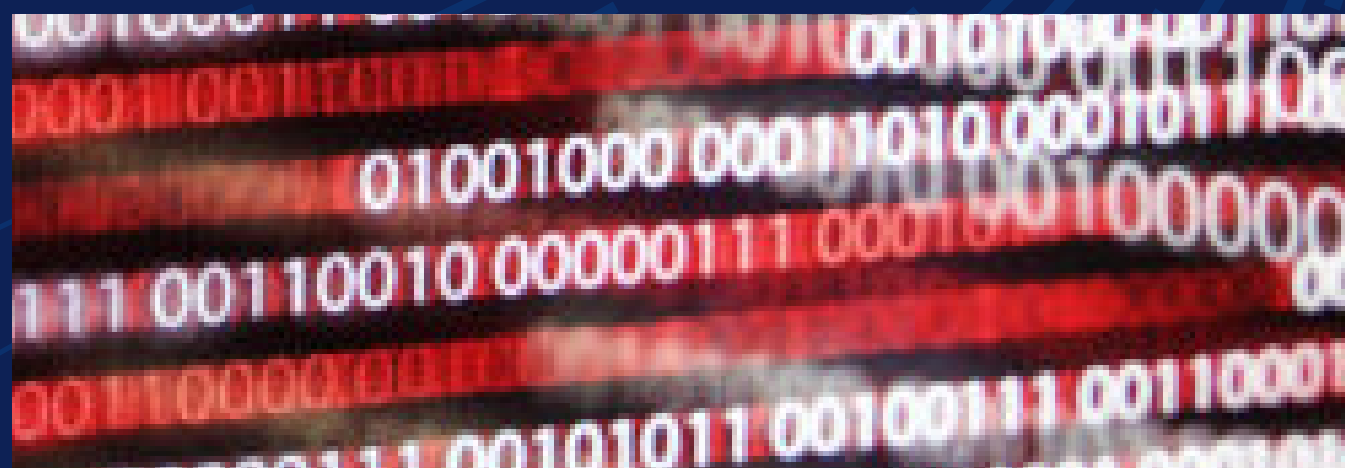


Doe het volgende op het gebied van cyberbeveiliging:

AI-beveiligingsmythes ontmaskeren



AI transformeert sectoren, maar als het gaat om het beveiligen van AI, worden veel organisaties het slachtoffer van mythes waardoor het complexer lijkt dan het echt is. De waarheid? Om AI-systemen te beschermen, is het niet nodig om helemaal vanaf nul te beginnen. Door bestaande cyberbeveiligingsprincipes op de unieke uitdagingen van AI toe te passen, bent u al een eind op weg.

Bij Dell Technologies begrijpen we de architectuur achter AI en kunnen we u helpen uw huidige oplossingen aan te passen aan dit nieuwe framework. Laten we de meest voorkomende mythes over AI-beveiliging ontrafelen en de waarheid onthullen om u te helpen uw systemen effectief te beveiligen.

Mythe 1: "AI-systemen zijn te complex om te beveiligen."

De waarheid: Het klopt dat AI nieuwe cyberbeveiligingsrisico's creëert, zoals promptinjectie, datamanipulatie en openbaarmaking van gevoelige informatie, om maar een paar voorbeelden te noemen. Bovendien hebben agentische AI-systemen ook een groter aanvalsoppervlak, aangezien ze kunnen worden misbruikt om resultaten te manipuleren of bevoegdheden te escaleren.

Hoewel het van cruciaal belang is deze kwetsbaarheden te herkennen en beveiligingsmaatregelen te implementeren om AI-systemen te beschermen tegen zowel traditionele als AI-specifieke bedreigingen, kunnen de risico's worden beheerd en AI-modellen worden beveiligd. Het is belangrijk om in gedachten te houden dat AI-systemen aanzienlijke hoeveelheden data nodig hebben als invoer en grote hoeveelheden data als uitvoer creëren. Dat stelt databescherming centraal als een van de belangrijkste beveiligingsstrategieën, samen met:

- Zero Trust-principes zoals identiteitsbeheer, op rollen gebaseerde toegang en continue verificatie.
- Regelmatige penetratietests en kwetsbaarheidsbeheer om zwakke punten te identificeren.
- Logboekregistratie en auditing om data-invoer en -uitvoer te valideren

Mythe 2: "Geen van mijn bestaande tools kan AI beveiligen."

De waarheid: AI-beveiliging draait niet om opnieuw beginnen, maar om slimmer werken met de tools die u al hebt. De meeste bestaande cyberbeveiligingstools kunnen worden aangepast om AI-systemen effectief te beveiligen. In de kern is AI een van de workloads in uw arsenaal die uw bedrijf aansturen, maar dan wel een met unieke kenmerken. Fundamentele cyberbeveiligingspraktijken, zoals identiteitsbeheer, netwerksegmentatie en -bewaking, eindpuntbescherming en databescherming, blijven essentieel voor het beschermen van AI-omgevingen. De sleutel is het aanpassen van deze praktijken om specifieke AI-uitdagingen aan te pakken, zoals het

beschermen van trainingsdata, het beveiligen van algoritmen en het beperken van risico's zoals vijandige input.

Een sterke verdediging begint met goede cyberhygiëne, zoals systeempatches, toegangscontrole en kwetsbaarheidsbeheer. Het is belangrijk om deze praktijken aan te passen, zodat u AI-specifieke risico's kunt aanpakken. Met AI-gerichte strategieën die zijn geïntegreerd in uw huidige beveiligingsaanpak en de juiste tools, wordt AI-beveiliging beheersbaar en effectief.

Het is echter belangrijk om te benadrukken dat bijgewerkte hardware een cruciale rol kan spelen bij het bestrijden van cyberaanvallen. Moderne AI PC's creëren bijvoorbeeld een sterke eerste verdedigingslinie tegen een grote aanvalsvector: eindpunten. Nu de ondersteuning voor Windows 10 eindigt, vormen verouderde pc's een risico. Bovendien vereist Windows 11 Trusted Platform Module (TPM) versie 2,0, een beveiligingschip die helpt bij versleuteling, veilig opstarten en bescherming tegen firmware-aanvallen. Veel oudere pc's hebben geen TPM of ondersteunen alleen een oudere versie. Dell biedt veilige commerciële AI PC's waarbij deze verbeteringen zijn geïntegreerd.

Hetzelfde geldt voor AI-infrastructuur zoals servers en storage. De Dell AI Factory omvat hardware die is geoptimaliseerd voor AI-beveiliging en een aantal ingebouwde beveiligingsfuncties bevat, variërend van een veilige leveringsketen en onveranderlijkheid van data tot isolatie en versleuteling.

Mythe 3: "AI-beveiliging draait alleen om het beschermen van data."

De waarheid: AI-beveiliging gaat verder dan fundamentele databescherming en omvat het beschermen van het hele AI-ecosysteem, inclusief modellen, API's, uitvoer, systemen en apparaten. Naarmate AI steeds meer wordt geïntegreerd in kritieke applicaties, nemen de risico's van misbruik toe. Zonder robuuste beveiligingsmaatregelen kan met AI-modellen worden geknoeid om schadelijke of misleidende uitvoer te genereren, kunnen API's worden misbruikt om onbevoegde toegang te krijgen tot gevoelige systemen en kan uitvoer onbedoeld privé- of vertrouwelijke informatie blootstellen.



Uitgebreide AI-beveiliging vereist een meerlaagse aanpak. Deze omvat het beschermen van modellen tegen aanvallen waarbij wordt geprobeerd invoerdata te manipuleren om AI-systemen te misleiden, het beveiligen van API's met sterke verificatiemethoden om onbevoegd gebruik te voorkomen, en het **continu bewaken van uitvoer** op ongebruikelijke of verdachte patronen die mogelijk wijzen op een aanval of storing. Effectieve AI-beveiliging garandeert niet alleen de integriteit en betrouwbaarheid van AI-systemen, maar bouwt ook vertrouwen op bij gebruikers en belanghebbenden door de risico's van kwaadwillig gebruik of onbedoelde gevolgen te beperken.

Mythe 4: "AI heeft geen menselijk toezicht nodig."

De waarheid: Governance en menselijk toezicht zijn van cruciaal belang om ervoor te zorgen dat AI-systemen ethisch, voorspelbaar en in

overeenstemming met menselijke waarden werken. Geavanceerde AI-systemen, met name agentische AI met autonome besluitvormingsmogelijkheden, introduceren unieke uitdagingen die robuuste beveiliging vereisen. Zonder het juiste toezicht kunnen deze systemen afwijken van de beoogde doelstellingen of onbedoeld gedrag vertonen dat risico's kan opleveren.

Om dit aan te pakken, is het essentieel duidelijke grenzen vast te stellen, gelaagde controlemechanismen te implementeren en te zorgen voor voortdurende menselijke betrokkenheid bij kritieke besluitvormingsprocessen. Regelmatige audits, transparantie in AI-activiteiten en grondige tests kunnen de verantwoordelijkheid en het vertrouwen verder verbeteren, misbruik helpen voorkomen en de verantwoorde implementatie van AI-technologieën bevorderen.

Best practices voor het versterken van AI-beveiliging

Om AI-specifieke beveiligingshiaten te dichten, moeten organisaties een proactieve en strategische aanpak hanteren. Hier zijn 10 best practices om uw AI-systemen te beveiligen:

- 

Gelaagde beveiligingsarchitectuur:
Gebruik segmentatie, firewalls en krachtige verificatie om uw infrastructuur, software en data op elke laag te beschermen.
- 

Beveilig de leveringsketen:
Implementeer een sterk leveranciersbeheerprogramma. Voer audits uit op leveranciers en componenten van derden, valideer integriteit en vertrouw op ondertekende code om kwetsbaarheden in de AI-ontwikkelingscyclus te voorkomen.
- 

Bescherm trainingsdata en -modellen:
Bescherm tegen vergiftigde data, vijandige invoer en andere bedreigingen door de integriteit van data te bewaken en robuuste validatietools toe te passen.
- 

Zorg voor strengere toegangscontroles:
Dwing principes met minimale bevoegdheden af, implementeer op rollen gebaseerde toegangscontrole (RBAC), verander aanmeldingsgegevens regelmatig en controleer machtigingen om onbevoegde toegang te voorkomen.
- 

Veilige API's:
Gebruik sterke verificatieprotocollen (zoals OAuth 2,0), dwing HTTPS-versleuteling af en werk API's regelmatig bij om mogelijke beveiligingslekken te dichten.
- 

AI-uitvoer bewaken en valideren:
Gebruik anomaliedetectie, logboekregistratie en waarschuwingen om te controleren op ongebruikelijke patronen of schadelijk gedrag in AI-uitvoer.
- 

Plan voor veerkracht:
Maak regelmatig een back-up van data en test plannen voor herstel na noodgevallen om downtime te minimaliseren en snel herstel te garanderen in geval van een inbreuk.
- 

Implementeer robuuste versleuteling:
Versleutel gevoelige data in rust en onderweg met behulp van sterke algoritmen, en beheer en vervang versleutelingssleutels regelmatig om de veiligheid te waarborgen.
- 

Voer regelmatig beveiligingsaudits en penetratietests uit:
Evalueer systemen regelmatig op kwetsbaarheden en gebruik penetratietests om risico's te ontdekken voordat ze kunnen worden misbruikt.
- 

Train personeel in best practices voor AI-beveiliging:
Train uw team regelmatig op het gebied van veilige ontwikkeling, bedreigingsherkenning en het onderhouden van sterke beveiligingspraktijken om inbreuken te voorkomen.



Het waardevoorstel van Dell: praktische AI-beveiligingsoplossingen.

AI-beveiliging lijkt misschien complex, maar is in werkelijkheid niet zo lastig. De waarheid? Het beveiligen van AI verschilt niet zo veel van het beveiligen van uw bestaande workloads: het gaat om het begrijpen van de architectuur en het toepassen van de juiste strategieën. Dat is precies waarbij Dell Technologies kan helpen.

We maken AI-beveiliging minder mysterieus door gebruik te maken van uw huidige oplossingen en ze naadloos te integreren in AI-gerichte architecturen. We pakken uitdagingen zoals promptinjectie,

API-misbruik en vijandige aanvallen aan zonder dat een volledige revisie van de infrastructuur nodig is.

De expertise van Dell ligt in het doorbreken van de mythes rond AI-beveiliging en het laten zien hoe haalbaar AI-beveiliging echt is. Of u nu net begint met uw AI-traject of uw verdediging wilt verbeteren, wij helpen u uw investeringen te beschermen, uw systemen te beveiligen en vol vertrouwen en effectief een veerkrachtige digitale toekomst op te bouwen. Laten we AI-beveiliging samen vereenvoudigen.

Dell producten en oplossingen die kunnen helpen

Aanbevolen Dell oplossing	Omschrijving
Dell AI Factory	Dell AI Factory beveiligt AI-workloads via een veilige leveringsketen, waardoor een betrouwbare infrastructuur wordt gegarandeerd, van ontwikkeling tot implementatie. Met functies zoals onveranderlijkheid, isolatie en versleuteling van data beschermt de oplossing gevoelige modellen en datasets, beschermt deze zich tegen cyberdreigingen en maakt deze schaalbare, efficiënte en naadloze AI-bewerkingen mogelijk in dynamische, datagestuurde omgevingen.
Cybertolerantie	PowerProtect beveiligt AI-workloads met geavanceerde functies zoals onveranderlijkheid en isolatie, waardoor data-integriteit en bescherming tegen cyberdreigingen worden gewaarborgd. De oplossing biedt end-to-end versleuteling en anomaliedetectie en maakt snel herstel mogelijk om downtime te minimaliseren.
Dell Trusted Workspace (eindpuntbeveiliging)	Een combinatie van geïntegreerde en optionele add-on-mogelijkheden die zijn ontworpen om commerciële AI PC's en AI-workloads die erop worden uitgevoerd te beveiligen. Ontwikkeld met veilige werkwijzen voor de leveringsketen en ingebouwde mogelijkheden zoals SafeBIOS en SafeID met TPM. Tot optionele add-ons behoren Secured Component Verification, SafeID met ControlVault en de partnersoftware CrowdStrike en Absolute voor maximale beveiliging van de werkruimte.
Adviesservices voor AI-beveiliging	Een reeks services die u kunnen helpen bij het ontwikkelen en implementeren van een uitgebreide AI-beveiligingsstrategie. Het aanbod omvat adviesservices, AI vCISO en planning voor databeveiliging.
Beheerde beveiligingsactiviteiten voor AI	Maakt diepgaande zichtbaarheid in de stack mogelijk om snel bedreigingen te detecteren en erop te reageren. Mogelijkheden omvatten Managed Detection and Response, Managed AI Guard, penetratietests voor AI en incidentrespons- en herstelservices.
Integratie van beveiligingssoftware	Ontwerp, installeer en configureer beveiligingstools die toegangsbeheer, applicaties, netwerken, clouds en meer beschermen.

Ontdek hoe u enkele van de grootste uitdagingen op het gebied van cyberbeveiliging van vandaag kunt aanpakken op dell.com/cybersecuritymonth