

Verantwoordelijke en ethische AI

Grondbeginselen voor het gebruik van
AI in Verint-producten

2025

VERINT
The CX Automation
Company™

Inhoudsopgave

Inleiding	3
Verint AI-strategie	3-4
Leidende principes voor AI	5
Passend: AI op de juiste manier gebruiken voor onze oplossingen	5
Veilig en privé: AI ontwerpen voor privacy, beveiliging en compliance	5
Eerlijk: AI bouwen en oneerlijke, gevoelige vooroordelen voorkomen	6
Verantwoording naar mensen: Betrouwbare output leveren die verantwoording aflegt aan mensen	6
Processen en controles	6
Ontwikkelings- en implementatiestandaarden	6
Verint AI-modelbeheer: regels voor het verkrijgen van data	7
AI-overwegingen en -richtlijnen	8-9
• Transparantie	
• Ethische besluitvorming	
• Vooringenomenheid en discriminatie	
• Privacy zorgen	
• Beveiligingsrisico's	
• Overwegingen op het gebied van wet- en regelgeving	
• Verkeerde informatie en manipulatie	
• Onbedoelde gevolgen	
Model Monitoring	10
• Passieve monitoring	
• Actieve monitoring	
Ondersteuning voor AI	10

Inleiding

Als leverancier van Customer Experience (CX) Automation-oplossingen waarin kunstmatige intelligentie (AI) is geïntegreerd, zet Verint zich in om ervoor te zorgen dat onze oplossingen niet alleen een aanzienlijke ROI opleveren voor onze klanten, maar ook AI op een ethische, verantwoorde manier benutten. Het Verint Open Platform is ontworpen met data en AI als kern, met de beste applicaties en AI-aangedreven bots die een breed scala aan functies automatiseren die van invloed zijn op de kosten en kwaliteit van CX. Dit document belicht onze AI-strategie, de principes die we volgen bij het uitvoeren van die strategie, en de processen, controles en richtlijnen die we hebben opgesteld om ons gebruik van AI te begeleiden.

Verint AI-strategie

Verint Da Vinci™ AI is de kern van het Verint Open Platform en is de drijvende kracht achter de applicaties die onze klanten elke dag gebruiken. Om gelijke tred te houden met de snelle innovatie in de AI-industrie, is de ontwikkeling van AI-modellen van Verint een open benadering waarbij de strategie zowel eigen modellen als modellen van derden omvat. Onze open benadering van AI betekent dat klanten kunnen vertrouwen op een investering in Verint Da Vinci AI voor de lange termijn.

Verint Da Vinci AI maximaliseert de impact door AI rechtstreeks in zakelijke workflows te injecteren, waardoor AI binnen handbereik van onze klanten komt. Door AI rechtstreeks in zakelijke workflows te injecteren, kunnen gebruikers van onze oplossingen:

- Ongestructureerde gegevens omzetten in intelligentie en actie.
- Machine learning gebruiken om bedrijfsprocessen te verbeteren.
- Anomalieën in gegevens detecteren en passende actie ondernemen.
- Trends en kansen identificeren met voorspellende modellering.



We hebben Verint Da Vinci AI ingebracht in een groeiend team van gespecialiseerde bots, die elk zijn gebouwd om AI te gebruiken, om één ding te doen en het goed te doen. Onze bots leven in het Verint Open Platform en worden ingezet voor gebruik zowel op het platform als met Verint-oplossingen die in andere clouds zijn geïmplementeerd. Klanten kunnen er over het algemeen voor kiezen om een of meer van onze bots in te zetten met hun Verint-oplossingen. Elke bot werkt aan het versterken van alle medewerkers van onze klanten en het verbeteren van CX-automatisering, bijvoorbeeld door:

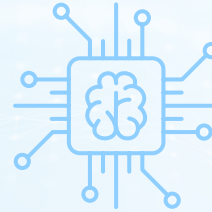
- Interacties samen te vatten;
- Het elimineren van werk na ieder gesprek;
- Reageren op vragen van klanten via selfservice en oproepen;
- Automatisch scoren van kwaliteitsevaluaties voor alle interacties;
- Voorspellen van klantinteracties over alle kanalen.



Onze AI-modellen zijn getraind op openbare, open source, gelicentieerde, synthetische, anonieme en geanonimiseerde datasets die zijn gegenereerd uit de Verint Engagement Data Hub om slimmere en snellere beslissingen te nemen. Deze getrainde AI-modellen zijn niet openbaar toegankelijk, er worden geen klantgegevens opgenomen in verbeteringen of nieuwe aanbiedingen, en deze modellen zijn alleen beschikbaar voor gebruik door onze klanten.

Engagement data kunnen bestaan uit:

- Interactiegegevens via verschillende kanalen
- Experience gegevens van zowel klanten als medewerkers.
- Workforce performance data



Het Verint Open Platform heeft unieke mogelijkheden om deze diverse data van meerdere applicaties samen te brengen tot een samenhangend geheel. De AI-modellen binnen Verint Da Vinci kunnen ook verder worden getraind op specifieke gegevens van de organisatie om hun vermogen om beter inzicht te bieden en CX-automatisering, te verbeteren. Wanneer een dergelijke aanvullende training plaatsvindt, zijn die uniek getrainde AI-modellen alleen beschikbaar voor gebruik door die specifieke organisatie.



De AI-services en -modellen van Verint Da Vinci worden gebouwd in Verint Labs, in combinatie met gegevens van de Engagement Data Hub binnen het Verint Open Platform. Verint Labs zijn veilige, gecontroleerde omgevingen die worden beheerd door een team van data- en AI-wetenschappers die actief samenwerken met klanten aan specifieke projecten, met de infrastructuur om tools en algoritmen van partners, externe onderzoeksorganisaties en universiteiten veilig te integreren. Dit team stuurt de onderzoeks-naar-productpijplijn van Verint Da Vinci en het Verint Open Platform aan door zich te richten op het verbeteren van onze AI- en machine learning-mogelijkheden.



Leidende principes voor AI

AI is een integraal onderdeel van het Verint Open Platform en is de drijvende kracht achter de applicaties die onze klanten dagelijks gebruiken. Verint realiseert zich dat geavanceerde technologieën zoals AI belangrijke uitdagingen kunnen opleveren die duidelijk, doordacht en direct moeten worden aangepakt.

De principes die ons gebruik van AI in ons aanbod sturen, zijn onder meer:

- AI alleen inzetten voor gepast gebruik.
- AI veilig en in overeenstemming met privacy- en andere regelgeving gebruiken.
- Werken om ervoor te zorgen dat AI op een eerlijke en veilige manier werkt.
- AI zo ontwerpen dat deze wordt gecontroleerd en verantwoording aflegt aan mensen.



Deze principes vormen de basis voor onze toewijding om AI-technologie op een verantwoorde en ethische manier te ontwikkelen. Ze zijn opgenomen in onze bedrijfscultuur en zijn ingebakken in de processen die worden gebruikt om de technologieën en oplossingen te leveren die we onze klanten bieden.

	Appropriate	Appropriate use of AI for our solutions
	Secure, Private & Compliant	AI is designed for privacy, security, and compliance
	Fair	Built to avoid unfair sensitive bias
	Accountable to People	Trusted output that is accountable to people

Passend: AI op de juiste manier gebruiken voor onze oplossingen

Verint werkt aan het beperken van potentieel schadelijke toepassingen bij het ontwerpen en implementeren van AI-technologieën. We beoordelen de AI-technologieën die we inzetten, hun gebruiksscenario's en de bijbehorende risico's. Verint evalueert ook hoe de AI-technologie zal worden geschaald over miljarden interacties en wereldwijde implementaties om ervoor te zorgen dat de AI-technologie die we aanbieden de waarde kan bieden die we voor ogen hebben.

Veilig, privé en compliant

AI-technologieën worden vaak blootgesteld aan grote hoeveelheden persoonlijke gegevens, wat problemen oproept met betrekking tot gegevensprivacy en -beveiliging. Verint ontwerpt en ontwikkelt AI-technologieën met behulp van best practices op het gebied van onderzoek en ontwikkeling. Verint werkt aan de ontwikkeling van zijn AI-systemen in overeenstemming met de toepasselijke wet- en regelgeving en ontwerpt dergelijke systemen op een manier die compliant gebruik door klanten mogelijk maakt. Dit omvat het respecteren van intellectuele eigendomsrechten van derden in het ontwerp- en ontwikkelingsproces. Verint ondersteunt ook de invoering van gegevensbeschermings-, nalevings- en privacyregelgeving om risico's in verband met AI aan te pakken, evenals de implementatie van veilige gegevensverwerkingspraktijken, en voert jaarlijkse beleidsevaluaties en updates uit om gelijke tred te houden met veranderingen.

Eerlijk: AI bouwen om oneerlijke, gevoelige vooroordelen te voorkomen

AI-algoritmen en datasets kunnen oneerlijke, gevoelige vooroordelen weerspiegelen en/of versterken. Verint begrijpt dat het onderscheiden van eerlijke en oneerlijke vooroordelen niet altijd eenvoudig is en kan verschillen tussen culturen en samenlevingen. We proberen te voorkomen dat het gebruik van onze AI-oplossingen leidt tot onrechtvaardige gevolgen voor mensen, met name die, welke verband houden met gevoelige kenmerken zoals ras, etniciteit, geslacht, nationaliteit, inkomen, seksuele geaardheid, handicap, leeftijd en politieke of religieuze overtuiging.

Verantwoording naar mensen: betrouwbare output leveren die verantwoording aflegt aan mensen

De AI-technologieën en gebruiksscenario's die Verint biedt, bieden passende mogelijkheden voor beoordelingen, relevante transparante uitleg en waarde. Onze AI-technologieën zijn onderworpen aan de juiste menselijke leiding en controle. Dit helpt bij het beheersen van risico's met betrekking tot ethische gebruiksscenario's, onbedoelde gevolgen en transparantie. Bovendien biedt het mechanismen om verkeerde informatie en manipulatie te helpen verminderen.

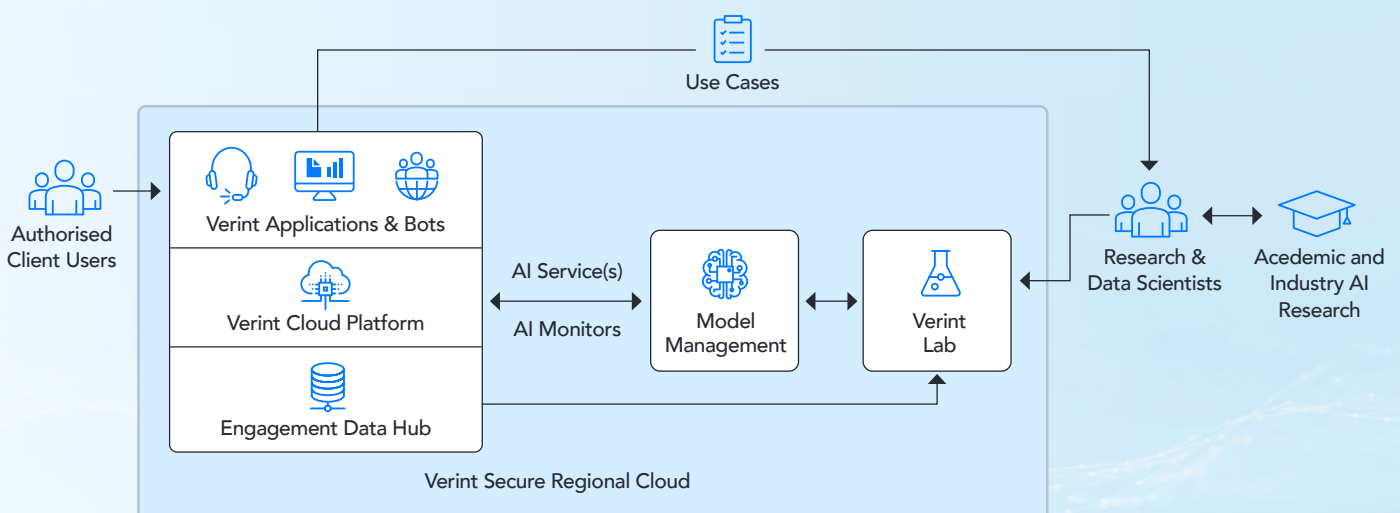
Processen en controles

Verint heeft standaardprocedures opgesteld en onderhoudt deze voor zijn onderzoeks-, engineering-, ontwikkelings- en operationele activiteiten wanneer AI in een oplossing wordt aangeboden, waaronder:

- Ontwikkelings- en implementatiewerkwijzen in overeenstemming met compliance standaarden.
- Modelbeheer en regels voor datagebruik en sourcing.
- Risicobeoordelingen ingebouwd in processen en beoordeeld door compliance-teams.
- Toezicht op het model.
- Ondersteuning via onze Verint Connect-fora.

Ontwikkelings- en implementatiestandaarden

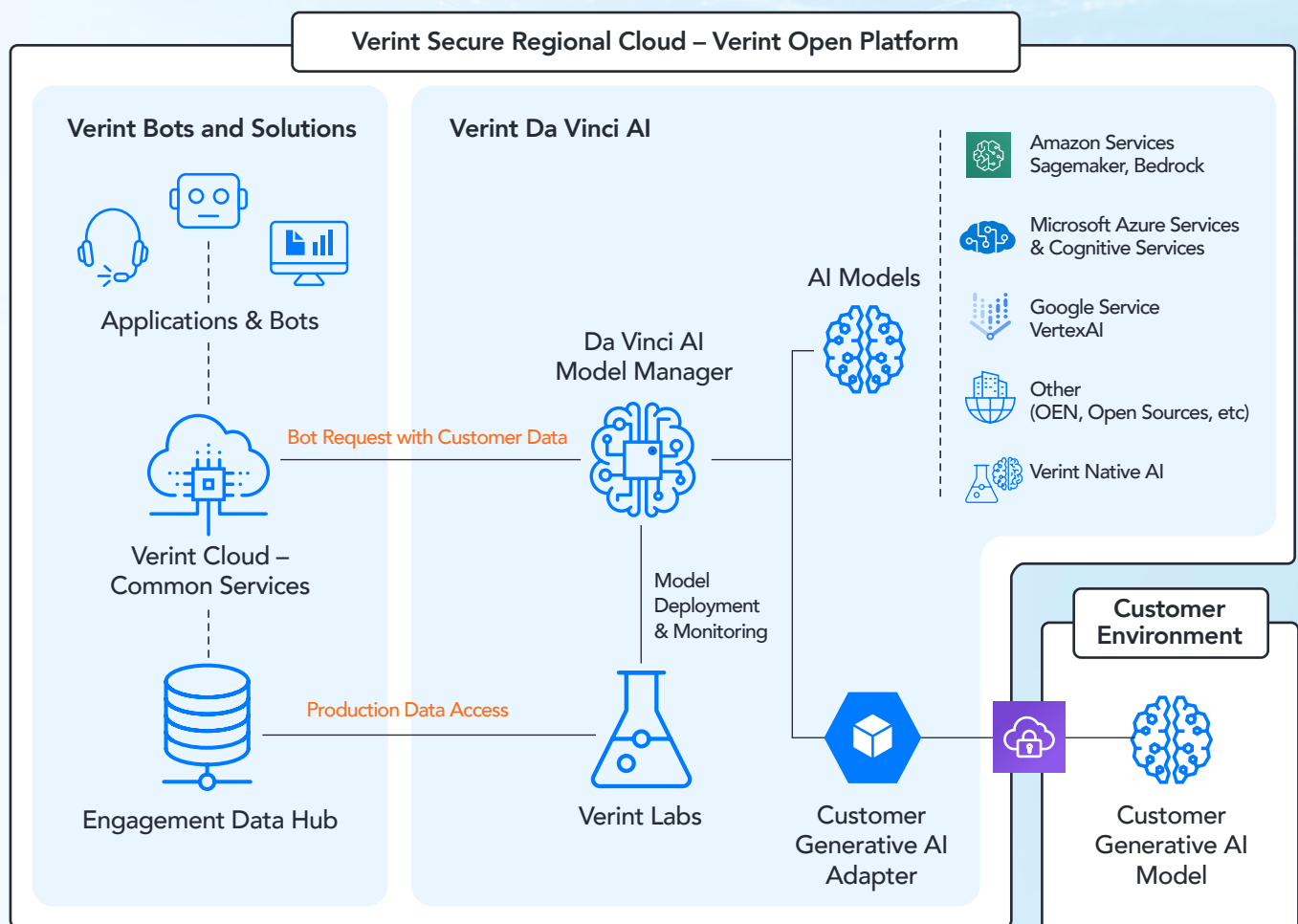
De AI-mogelijkheden van Verint zijn ontwikkeld en bevinden zich binnen het Verint Open Platform. Het Open Platform van Verint omvat veilige regionale cloudomgevingen waarin Verint de industriestandaard voor beveiliging handhaaft (<https://www.verint.com/verint-cloud-security-and-compliance/>) die de richtlijnen volgen die zijn uiteengezet door bijv. PCI, GDPR, ISO, SOC2 en HIPAA. Verint heeft passende technische en organisatorische maatregelen geïmplementeerd en handhaaft deze om persoonsgegevens te beschermen tegen onopzettelijke, ongeoorloofde of onwettige toegang, openbaarmaking, wijziging, verlies of vernietiging. Onze beveiligingspraktijken zijn van toepassing op activa en data die onder onze controle vallen, evenals op onze operationele omgevingen.



Verint AI Model Management: Data Use Sourcing Rules

Verint heeft ook praktijken vastgesteld voor het bouwen, implementeren en beheren van AI-modellen en beperkingen voor datagebruik:

- Klantgegevens verlaten de beveiligde regionale cloudomgeving van Verint niet bij het gebruik van onze AI-services.
- Verint kan klantgegevens openen, verwerken en gebruiken bij het verbeteren of creëren van verbeteringen of nieuwe aanbiedingen met betrekking tot de SaaS-service, maar er worden geen klantgegevens daadwerkelijk opgenomen in dergelijke verbeteringen of nieuwe aanbiedingen.
- De AI-bepaling van Verint in haar Master Service Agreement (MSA) maakt duidelijk:
 - Inputs en outputs zijn vertrouwelijke informatie van onze klant
 - Zonder schriftelijke toestemming van een klant zal Verint AI-modellen niet rechtstreeks trainen met klantgegevens



Houd er rekening mee dat AI-services van derden werken binnen de Verint Secure Regional Cloud-omgevingen. Deze services worden uitgevoerd in een privé, niet-openbare, configuratie. De services worden volledig beheerd door Verint en staan niet toe dat die AI-modellen van derden worden getraind met behulp van klantgegevens, en/of bieden geen toegang tot input en output van die derden.

AI-overwegingen en -richtlijnen

Verint heeft een framework met overwegingen en richtlijnen opgesteld voor de belangrijkste factoren die een rol spelen bij de ontwikkeling, implementatie en het gebruik van AI-technologieën. Het framework gaat in op enkele veelgestelde vragen over de impact van AI op deze oplossingen. Is de dienst bijvoorbeeld alleen bedoeld om informatie te verstrekken waarop mensen kunnen reageren, of is het bedoeld voor het nemen van geautomatiseerde beslissingen? Hoe verklaarbaar zijn de voorspellingen van het model? Biedt de implementatie inzicht zodat we kunnen begrijpen waarom de engine de voorspelling deed die hij deed? Wordt AI gebruikt op een manier die van invloed kan zijn op de menselijke werkgelegenheid of het loon?

Verint probeert deze overwegingen aan te pakken in ons oplossingsaanbod, zoals hierboven besproken, en evalueert ze vanuit de volgende perspectieven:



Transparantie

Gebrek aan transparantie in AI-systemen, met name in deep learning-modellen, is een dringend probleem. Deze ondoorzichtigheid verhuult de besluitvormingsprocessen en de onderliggende logica van deze technologieën. Om inzicht te geven in die ondoorzichtige gebieden, proberen we de context en reikwijdte te verstrekken van de informatie die de onderliggende algoritmen en processen gebruiken. Dit omvat:

- **Displaytransparantie:** Biedt een real-time inzicht in de taak die door AI wordt uitgevoerd en in de acties die het AI-systeem onderneemt.
- **Verklaarbaarheid:** Biedt op een terugkijkende manier informatie over de logica, het proces, de factoren of de redenering waarop de acties of aanbevelingen van het systeem zijn gebaseerd.

Ethische besluitvorming

Het bijbrengen van ethische waarden in AI-systemen, vooral in besluitvormingscontexten met aanzienlijke gevolgen, vormt een aanzienlijke uitdaging. Onderzoekers en ontwikkelaars moeten prioriteit geven aan de ethische implicaties van AI-technologieën om negatieve maatschappelijke gevolgen te voorkomen.

We proberen dit aan te pakken door:

- Het uitvoeren van een ethische beoordeling van AI voor nieuwe projecten als onderdeel van het onderzoek en definitieproces. Het ethische AI-beoordelingsproces onderzoekt de mogelijke schade die kan worden veroorzaakt door normaal gebruik van de AI: met resultaten van hoge kwaliteit; met onjuiste resultaten of resultaten van lage kwaliteit; hoe de technologie verkeerd kan worden toegepast; en hoe de AI correct of onjuist kan leren als het wordt gebruikt. De ethische beoordeling van Verint omvat de volgende vragen en een bespreking van hoe elk probleem kan worden beperkt:
- Welke mogelijke schade ontstaat er wanneer de technologie wordt gebruikt zoals bedoeld en correct functioneert?

- Welke mogelijke schade ontstaat er wanneer de technologie wordt gebruikt zoals bedoeld, maar onjuiste resultaten geeft?
- Welke mogelijke schade ontstaat er als gevolg van mogelijk misbruik van de technologie?
- Als het systeem leert van gebruikersinvoer zodra het is geïmplementeerd:
 - Welke controles en beperkingen zijn bedoeld met betrekking tot het leren van het systeem?
 - Bevat dat leren klantspecifieke gegevens?
 - Is het geleerde van toepassing op het aanbod in het algemeen, of alleen op tenant level?

Vooringenomenheid en discriminatie

AI-systemen kunnen onbedoeld maatschappelijk gevoelige vooroordelen in stand houden of versterken als gevolg van bevooroordeelde trainingsgegevens of algoritmisch ontwerp. Om discriminatie tot een minimum te beperken en eerlijkheid te waarborgen, is het van cruciaal belang om te investeren in de ontwikkeling van onbevooroordeelde algoritmen en diverse trainingsdatasets.

We proberen dit aan te pakken door:

- Bewust zijn van algemeen bekende vooroordelen die aanwezig kunnen zijn in AI-systemen, zoals databias, algoritmische bias en bevestigingsbias.
- Periodiek beoordelen, evalueren en aanpakken van door AI gegenereerde output tijdens het ontwikkelingsproces op mogelijke vooroordelen en onnauwkeurigheden.
- Het aanpakken van bias- of nauwkeurigheidsproblemen die via ons Incident Management System worden gemeld als een bug/defect.
- Het gebruik van AI-systemen die door leveranciers worden geleverd met transparante methodologieën en documentatie om hun besluitvormingsprocessen beter te begrijpen.
- Documenteren en communiceren van eventuele geïdentificeerde vooroordelen en mitigatie-inspanningen.

Privacy zorgen

AI-technologieën verzamelen en analyseren vaak grote hoeveelheden persoonlijke gegevens, wat problemen oproept met betrekking tot de privacy en beveiliging van gegevens. Om privacyrisico's te beperken, streven we ernaar ons te houden aan veilige standaarden voor gegevensverwerking.

We proberen dit aan te pakken door:

- Te voldoen aan privacy- en andere wetten die van toepassing zijn op Verint als leverancier van AI.
- Het naleven van de Gedragscode, Global Information Security Policy en het privacybeleid van Verint bij de toegang tot en verwerking van persoonsgegevens.
- Het nakomen van onze contractuele verplichtingen met betrekking tot de omgang met klantgegevens, ook met betrekking tot persoonsgegevens.
- Het gebruik van veilige en aan gegevensprivacy gekoppelde omgevingen, zoals Verint Labs, om AI-systemen of -services te onderzoeken, te ontwikkelen en te testen.

Beveiligingsrisico's

Naarmate AI-technologieën steeds geavanceerder worden, nemen ook de beveiligingsrisico's die gepaard gaan met het gebruik ervan en de kans op misbruik toe. Hackers en kwaadwillenden kunnen de kracht van AI benutten om geavanceerdere cyberaanvallen te ontwikkelen, beveiligingsmaatregelen te omzeilen en kwetsbaarheden in systemen te misbruiken.

We proberen dit aan te pakken door:

- Ons te houden aan het interne AI-gebruiksbeleid van Verint, deze richtlijnen en onze Global Information Security Policy bij het gebruik en de ontwikkeling van AI-systemen.
- Het gebruik van veilige en aan gegevensprivacy gekoppelde omgevingen, zoals Verint Labs, om AI-systemen of -services te onderzoeken, te ontwikkelen en te testen.

Overwegingen op het gebied van wet- en regelgeving

Over de hele wereld worden nieuwe wettelijke kaders en voorschriften aangenomen om de unieke problemen aan te pakken die voortvloeien uit AI-technologieën, waaronder toewijzing van verantwoordelijkheden, aansprakelijkheden en intellectuele eigendomsrechten. Rechtssystemen evolueren om gelijke tred te houden met de technologische vooruitgang en de rechten van individuen en bedrijven te beschermen.

De rechten en wetten die van toepassing kunnen zijn op aanbiedingen van Verint zijn onder meer:

- Wetgeving inzake gegevensbescherming en privacy.
- Intellectueel eigendomsrecht
- Wetten die van toepassing zijn op werkgelegenheid, de bescherming van beschermde groepen of andere wet- en regelgeving die verband houdt met belangen voor mensen.
- Wetten die het gebruik van AI zelf reguleren.

We proberen dit aan te pakken door:

- AI-systemen te ontwikkelen die zijn ontworpen om te voldoen aan de toepasselijke wet- en regelgeving en dergelijke systemen zo te ontwerpen dat klanten ze kunnen gebruiken in overeenstemming met de regelgeving.
- Het respecteren van intellectuele eigendomsrechten van derden in het ontwerp- en ontwikkelingsproces.
- Het naleven van het interne AI-gebruiksbeleid van Verint en deze richtlijnen bij het ontwikkelen van AI-services of -systemen die door onze klanten worden gebruikt.
- Monitoren en beoordelen van toepasselijke updates van wet- en regelgeving en best practices.

Verkeerde informatie en manipulatie

Door AI gegenereerde inhoud, zoals deepfakes, draagt bij aan de verspreiding van valse informatie en de manipulatie van de publieke opinie. Inspanningen om door AI gegenereerde desinformatie op te sporen en te bestrijden zijn van cruciaal belang voor het behoud van de integriteit van informatie in het digitale tijdperk. Gegevens van individuen (zoals klanten van onze klanten) kunnen worden verzameld en ingevoerd in AI-services voor verdere verrijking of automatisering en omvatten vaak het soort informatie (d.w.z. persoonlijke gevoelige informatie, spraakaudio-opnamen, foto's, accountinformatie, enz.) die door kwaadwillende actoren kunnen worden gebruikt om een persoon valselyk voor te stellen. Het onder de loep nemen van de beveiliging van elke AI-service die we gebruiken, en ervoor zorgen dat de AI-service niet optreedt namens of zichzelf voordoeft als de consument, zijn kernaandachtspunten voor ons.

We proberen ze aan te pakken door:

- Ervoor te zorgen dat de gegevensinvoer die naar een AI-service gaat, wordt beschermd, en het veilig waarborgen om onderschepping of gebruik door derden voor alternatieve doeleinden te voorkomen.
- Ervoor zorgen dat de AI-dienst niet optreedt namens of zichzelf voordoeft als de consument.

Onbedoelde gevolgen

AI-systemen kunnen vanwege hun complexiteit en gebrek aan constant menselijk toezicht onverwacht gedrag vertonen of beslissingen nemen met onvoorziene gevolgen. Deze onvoorspelbaarheid kan leiden tot resultaten die een negatieve invloed hebben op individuen, bedrijven of de samenleving. Robuuste test-, validatie- en monitoringprocessen kunnen ontwikkelaars en onderzoekers helpen dit soort problemen te identificeren en op te lossen voordat ze escaleren.

We proberen dit aan te pakken door:

- AI-modellen te ontwerpen die voor onze diensten worden gebruikt die uitgelegd kunnen worden aan de menselijke beslisser.
- Het opnemen in de productdocumentatie van beschrijvingen van de AI-modellen die worden gebruikt met het Verint Open Platform.
- Het opzetten van een goedkeuringsproces voor ontwikkeling dat vereist dat wordt beoordeeld hoe AI-mogelijkheden worden gecombineerd met andere Verint-functionaliteit, en hoe de AI zal worden toegepast binnen het Verint-aanbod om verdere ontwikkelingsgebruiksscenario's zonder vereiste goedkeuringen te beperken.

Model Monitoring

Het Verint Open Platform biedt mogelijkheden om zowel actief (d.w.z. voortdurend observeren en analyseren) als passief (d.w.z. problemen die worden gemeld via onze standaard supporttickets) AI-services te monitoren die zijn ingezet voor klanten in regionale clouds over de hele wereld. Modelmonitoring dient als een maatregel om de gezondheid, efficiëntie en effectiviteit van AI-toepassingen of bots te behouden.

Passieve monitoring

Verint biedt een mechanisme voor klanten, partners en medewerkers van Verint (d.w.z. leden van het professionele serviceteam die namens klanten werken) om ondersteuningstickets in te dienen over applicatie- of botgedrag dat onverwacht is of moet worden onderzocht. De serviceorganisatie van Verint zal achter de schermen coördineren met de juiste product-, engineering- en onderzoeksteams om het probleem aan te pakken of een vraag te beantwoorden.

Actieve monitoring

Actieve modelmonitors kunnen worden ingezet als onderdeel van de projectdefinitie binnen engineering/architectuur, onderzoek en ontwikkeling, afhankelijk van het gebruiksscenario. Elke toepassing of bot die AI implementeert, kan de prestaties bewaken en verbeteringen opnemen in het kernmodel, aanvullende trainingsgegevens verstrekken, aanvullende modeltrainingsgebeurtenissen uitvoeren of andere herstelltaken implementeren om problemen aan te pakken die door de monitor zijn geïdentificeerd.

Ondersteuning voor AI

Verint Connect is het Verint-extranet voor klanten en partners en is het primaire Verint-contactpunt voor deze partijen om vragen te stellen, informatie te vinden, problemen te melden en trainingsmateriaal te ontvangen. Als er een vraag, probleem of probleem met een AI-component van een Verint-oplossing wordt aangetroffen, kan dit worden gemeld via Verint Connect met het juiste prioriteitsniveau voor een reactie door Verint.



Verint®. The CX Automation Company™

Europe, Middle East & Africa

info.emea@verint.com

+44(0) 1932 839500

Americas

info@verint.com

+1 770 754 1900

1-800-4VERINT

Asia Pacific

info.apac@verint.com

+ (852) 2797 5678



verint.com



x.com/verint



linkedin.com/company/verint



verint.com/blog