



**Voorbeeldexamen**

Editie 202603

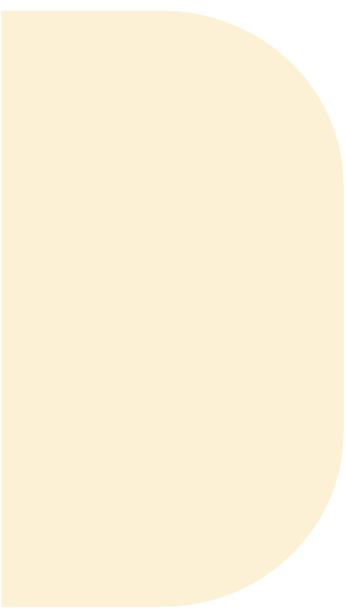
Copyright © EXIN Holding B.V. 2026. All rights reserved.  
EXIN® is a registered trademark.

No part of this publication may be reproduced, stored, utilized or transmitted in any form or by any means, electronic, mechanical, or otherwise, without the prior written permission from EXIN.



# Inhoud

|                 |    |
|-----------------|----|
| Inleiding       | 4  |
| Voorbeeldexamen | 5  |
| Antwoordsleutel | 22 |
| Evaluatie       | 58 |



# Inleiding

Dit is het EXIN Artificial Intelligence Compliance Professional (AICP.NL) voorbeeldexamen. Op dit examen is het Reglement voor de Examens van EXIN van toepassing.

Dit examen bestaat uit 40 meerkeuzevragen. Elke vraag heeft een aantal antwoorden, waarvan er één correct is.

Het maximaal aantal te behalen punten is 40. Elke goed beantwoorde vraag levert u 1 punt op. U hebt minimaal 26 punten nodig om te slagen.

De beschikbare tijd is 90 minuten.

U mag AI Act gebruiken tijdens dit examen.

Veel succes!

# Voorbeeldexamen

1 / 40

De AI Act is een wet die is opgesteld voor de Europese Unie (EU). In artikel 1 worden de doelstellingen van de AI Act beschreven.

Wat zijn de **belangrijkste** doelstellingen van de AI Act?

- A) Richtlijnen die puur bedoeld zijn voor milieubescherming, zonder specifieke regels voor AI met een hoog risico, zonder verboden, en met innovatiemaatregelen alleen voor grote bedrijven
- B) Geharmoniseerde regels voor AI-systemen in de EU, verboden voor bepaalde AI-praktijken, vereisten voor AI met een hoog risico, transparantieregels, markttoezicht en innovatieondersteuning
- C) Verboden voor AI-praktijken, uitsluitend regels voor AI voor algemene doeleinden, transparantieregels behalve voor AI met een hoog risico en innovatieondersteuning beperkt tot niet-Europese entiteiten
- D) Regels voor AI-systemen beperkt tot veiligheid en gezondheid, verboden voor alle AI-praktijken, transparantieregels uitsluitend voor AI met een hoog risico, en innovatieondersteuning voor alle bedrijven behalve start-ups

2 / 40

Wat betekenen verantwoording en naleving volgens de AI Act?

- A) Verantwoording draait om het handhaven van gebruikersprivacy en dataveiligheid, terwijl naleving betrekking heeft op integratie met een bestaande IT-infrastructuur.
- B) Verantwoording houdt in dat ontwikkelaars en operators die AI ontwikkelen daarvoor verantwoordelijk worden gehouden, en naleving betekent dat wettelijke vereisten in acht worden genomen.
- C) Verantwoording draait om het garanderen dat AI-systemen winstgevend zijn voor ontwikkelaars, en naleving houdt in dat systemen voldoen aan de wensen en voorkeuren van gebruikers.
- D) Verantwoording betekent dat AI-gebruikers verantwoordelijk worden gehouden voor correct gebruik van het systeem, terwijl naleving betekent dat binnen de sector gangbare normen voor AI-innovatie in acht worden genomen.

3 / 40

Personen die de gevolgen ondervinden van AI-systemen hebben onder de AI Act specifieke rechten om transparantie, billijkheid en verantwoording te waarborgen.

Welk recht wordt expliciet verleend onder de AI Act?

- A) Het recht om te worden geïnformeerd over het feit dat ze interageren met of de gevolgen ondervinden van een AI-systeem
- B) Het recht om inzage in de broncode van het AI-systeem te eisen
- C) Het recht om het gebruik van AI te verbieden voor elk besluitvormingsproces dat betrekking op hen heeft
- D) Het recht om verwijdering van door het AI-systeem gebruikte persoonsgegevens aan te vragen

#### 4 / 40

Anna is als compliance officer bij een kleine of middelgrote onderneming (kmo) verantwoordelijk voor het toezicht op de implementatie van een nieuw AI-systeem dat zal worden ingezet om de klantondersteuning te automatiseren. Het bedrijf heeft dit systeem niet zelf gebouwd, maar koopt het van een andere aanbieder. Onder de AI Act wordt dit systeem geclassificeerd als een AI-systeem met een hoog risico.

Anna is gevraagd om te zorgen dat het bedrijf bij de inzet en monitoring van dit AI-systeem voldoet aan de verplichtingen van de gebruiker. Ze moet bepalen welke acties prioriteit moeten krijgen en welke acties moeten worden vermeden.

Wat kan Anna beter **niet** overwegen, gezien de verplichtingen voor AI-gebruikers?

- A) De algoritmes van het AI-model verder ontwikkelen om de besluitvormende capaciteiten te verbeteren zonder de aanbieder daarbij te betrekken
- B) De prestaties van het AI-systeem gedetailleerd bijhouden en zorgen dat aan alle relevante rapportagevereisten wordt voldaan
- C) De prestaties van het AI-systeem monitoren om te waarborgen dat het werkt zoals beoogd en voldoet aan veiligheidsnormen
- D) Elk ernstig incident met of gebrekkig functioneren van het AI-systeem melden bij de toepasselijke instanties, zoals wettelijk vereist

#### 5 / 40

Een AI-systeem voor gezichtsherkenning wordt in openbare ruimtes gebruikt voor beveiligingsdoeleinden. Eén organisatie is voor dit AI-systeem het meest relevant voor toezicht op de naleving van alle regelgeving op het gebied van databescherming en privacy (zoals de Algemene Verordening Gegevensbescherming (AVG)).

Welke organisatie is dat?

- A) De Europese Consumentenorganisatie (Bureau Européen des Unions de Consommateurs, BEUC)
- B) De Europese Raad voor Artificiële Intelligentie (European Artificial Intelligence Board; AI-board)
- C) Het Europese Hof van Justitie (EHJ)
- D) Het Europese Comité voor Gegevensbescherming (EDPB)

**6 / 40**

Een bedrijf ontwikkelt een AI-systeem voor gepersonaliseerde marketing. Dit systeem maakt gebruik van machine learning algoritmes (ML-algoritmes) om advertenties af te stemmen op individuele klanten. Tijdens een nalevingsbeoordeling stelt het team de volgende risico's vast:

- Er is geen documentatie waaruit duidelijk blijkt hoe het AI-systeem data verwerkt.
- Het proces aan de hand waarvan het AI-systeem gepersonaliseerde aanbevelingen doet, wordt niet volledig begrepen.
- Klanten klagen over deze kwesties.

Het bedrijf moet de AI Act naleven en hanteert daarvoor de norm ISO/IEC 42001 en het NIST AI Risk Management Framework (RMF).

Wat moet het bedrijf volgens deze norm en dit kader doen om deze kwesties aan te pakken?

- A) Een reeks tests van de gebruikerservaring (UX) uitvoeren om feedback over de bruikbaarheid, leerbaarheid en klantvoorkeuren te verzamelen
- B) Proberen om voorspellingen van het systeem nauwkeuriger te maken om zo de kostenefficiëntie, klanttevredenheid en betrokkenheid te verbeteren
- C) Een documentatieproces implementeren waarin databronnen, verwerkingsmethoden en de algoritmische besluitvorming uitvoerig worden beschreven
- D) De systeemhardware upgraden om snellere verwerking, meer efficiëntie en een hogere klanttevredenheid te waarborgen

**7 / 40**

Een bedrijf ontwikkelt een AI-systeem voor de monitoring van patiënten die in een ziekenhuis worden opgenomen. Het systeem werkt met high-definition camera's (HD-camera's) op de kamers om de status van patiënten in real time te monitoren. Als het systeem detecteert dat een patiënt onrustig is, wordt er automatisch iemand van de verpleging opgeroepen om langs te gaan.

Om de prestaties van het AI-systeem te verbeteren, wil het bedrijf een videodatabase van patiënten aanleggen, waarbij video's op kritieke punten worden voorzien van een opmerking van een professional. Zo komt er meer trainingsdata voor het systeem beschikbaar.

Het bedrijf overweegt een gegevensbeschermingseffectbeoordeling (DPIA) uit te voeren. Het verantwoordelijke team twijfelt of een DPIA nodig is. Als een DPIA verplicht is, wil het team weten wanneer de beoordeling dient te worden uitgevoerd: nu of pas nadat de update is uitgerold.

Het bedrijf moet de AI Act en de Algemene Verordening Gegevensbescherming (AVG) naleven.

Moet het bedrijf nu al een DPIA uitvoeren?

- A) Ja, want een DPIA is verplicht voor AI-projecten die mogelijk een hoog risico vormen voor de rechten van natuurlijke personen.
- B) Ja, want een DPIA is verplicht voor elk project waarbij persoonsgegevens worden verzameld, zelfs als het risico van het project laag is.
- C) Nee, want een DPIA is niet verplicht wanneer gegevens worden gebruikt voor trainingsdoeleinden, onderwijs of wetenschappelijk onderzoek.
- D) Nee, want een DPIA is pas verplicht nadat een AI-systeem volledig is ontwikkeld, getest en uitgerold.

**8 / 40**

Een bedrijf ontwikkelt een AI-systeem voor gezichtsherkenning in real time. Een particulier beveiligingsbedrijf zet het AI-systeem in voor de monitoring van een openbaar overdekt winkelcentrum. Het systeem scant alle bezoekers, vergelijkt hen met een database met mensen die eerder een overtreding hebben begaan en een database met politiek activisten, en markeert bezoekers die voorkomen in een van beide databases. Gemarkeerde bezoekers worden gedurende hun hele bezoek onopgemerkt gevolgd om te beoordelen of zij zich volgens het beveiligingsbedrijf verdacht gedragen.

In welke categorie moet het gebruik van dit AI-systeem volgens de AI Act worden ingedeeld?

- A) Onaanvaardbaar risico
- B) Hoog risico
- C) Beperkt risico
- D) Laag of geen risico

**9 / 40**

Een reisbureau gebruikt een AI-systeem om dynamische, gerichte marketingcampagnes voor pakketreizen te ontwikkelen. De campagnes omvatten realtimeadvertenties op social media en reisplatforms, waarbij gebruik wordt gemaakt van de browsegeschiedenis van sitebezoekers. Het reisbureau gebruikt AI om de gemoedstoestand van gebruikers te bepalen en stemt vervolgens de suggesties voor bestemmingen en activiteiten hierop af.

Het reisbureau moet de AI Act naleven.

Welk risico moet het reisbureau aanpakken?

- A) Het risico op mogelijke bias. Het bureau moet de trainingsdata regelmatig bijwerken om te voorkomen dat er suggesties voor irrelevante bestemmingen verschijnen of een onjuiste gemoedstoestand wordt bepaald.
- B) Het risico op ineffectieve reclameactiviteiten. Het bureau moet vooral proberen het algoritme bij te werken, want gepersonaliseerde advertenties vallen niet onder de AI Act.
- C) Het risico op een gebrek aan transparantie. Het bureau moet openheid over de AI garanderen, bias in suggesties terugdringen en beoordelen of de reclameactiviteiten ethisch zijn.
- D) Het risico op misbruik van persoonsgegevens. Het bureau moet stoppen met het gebruik van AI-gestuurde personalisatie omdat het onder de AI Act verboden is om persoonsgegevens te gebruiken voor gerichte advertenties.

**10 / 40**

Een bedrijf heeft een AI-model ontwikkeld dat in verschillende sectoren kan worden gebruikt, waaronder de gezondheidszorg en financiële dienstverlening. Vanwege de brede toepassingsmogelijkheden brengt het AI-model mogelijk risico's voor de volksgezondheid met zich mee.

Wat heeft het bedrijf ontwikkeld en welke praktijken dient het bedrijf volgens de AI Act te implementeren?

- A) Het bedrijf heeft een AI-systeem voor algemene doeleinden ontwikkeld dat systeemrisico's met zich meebrengt. Het dient aanvullende tests uit te voeren om die risico's te beperken.
- B) Het bedrijf heeft een AI-systeem met een hoog risico ontwikkeld. Het dient alle vereisten voor AI-systemen met een hoog risico te implementeren zoals beschreven in de AI Act.
- C) Het bedrijf heeft een beperkt (narrow) AI-model ontwikkeld. Om risico's te voorkomen dient het te waarborgen dat het model alleen werkt binnen vooraf gedefinieerde parameters.
- D) Het bedrijf heeft een experimenteel AI-model ontwikkeld. Het dient zich op onderzoek en ontwikkeling te richten zonder onmiddellijk risicomanagement.

**11 / 40**

Een organisatie ontwikkelt een AI-systeem met een hoog risico. Tijdens het testen stelt het ontwikkelteam diverse risico's vast, waaronder inconsistenties in de volledigheid van data en de aanwezigheid van verouderde records. Deze risico's hebben mogelijk een negatieve impact op de prestaties van het model.

De organisatie moet de AI Act naleven en gebruikt daarvoor het kader CEN/CLC/TR 18115.

Wat moet de organisatie volgens dit kader doen om deze risico's aan te pakken?

- A) Een gegevensbeschermingseffectbeoordeling (DPIA) uitvoeren om de billijkheid van besluitvorming door AI aan te pakken
- B) Alle datasets voor trainings- en testdoeleinden versleutelen aan de hand van protocollen om onbevoegde inzage in persoonsgegevens te voorkomen
- C) Algemene maatregelen voor risicobeheer implementeren om de genoemde operationele risico's en reputatierisico's te beperken
- D) De datakwaliteit verbeteren door gestructureerde meetwaarden voor de kwaliteit en statistische evaluatiemethoden toe te passen

**12 / 40**

In de AI Act worden diverse rollen gerelateerd aan een AI-systeem beschreven.

Wat is de definitie van de rol 'importeur van een AI-systeem'?

- A) Een persoon of organisatie die een AI-systeem ontwerpt, ontwikkelt en onder diens eigen naam of merk in de handel brengt.
- B) Een persoon of organisatie die een AI-systeem in de handel brengt, maar niet verantwoordelijk is voor de oorspronkelijke ontwikkeling ervan.
- C) Een persoon of organisatie die een AI-systeem gebruikt voor eigen activiteiten en zorgt dat alle lokale verplichtingen van de gebruiker worden nageleefd.
- D) Een regelgevende instantie die als taak heeft te monitoren of het AI-systeem wordt geïmporteerd conform de bepalingen van de AI Act.

**13 / 40**

Een bedrijf ontwikkelt een AI-systeem voor fraudedetectie bij financiële transacties. Dit systeem analyseert transactiepatronen om verdachte activiteiten op te sporen en frauduleus gedrag te voorkomen. Gezien de kans op valspositieven die mogelijk een nadelige invloed op legitieme transacties hebben en het feit dat fraudetactieken voortdurend veranderen, erkent het bedrijf dat effectieve waarborgen vereist zijn.

Het bedrijf moet de AI Act naleven. Om problemen met valspositieven te helpen voorkomen, hanteert het bedrijf de norm ISO/IEC 23894.

Wat moet het bedrijf volgens deze norm doen om deze problemen te voorkomen?

- A)** Risicomanagement integreren in alle activiteiten om breed toezicht en proactieve risicobeperking te waarborgen
- B)** Maatregelen voor dataprivacy uitbreiden om gevoelige informatie te beschermen en te voldoen aan privacyregels
- C)** Proberen de nauwkeurigheid van het model te verbeteren om betrouwbare prestaties te waarborgen en valspositieven tot een minimum te beperken
- D)** Cyberbeveiligingsmaatregelen implementeren om het systeem te beschermen tegen dreigingen van buitenaf en onbevoegde inzage

**14 / 40**

Een productiebedrijf maakt in zijn assemblagelijnen gebruik van AI-gestuurde robotsystemen voor kwaliteitscontrole. Het onderzoeksteam verneemt dat een anonieme klokkenluider meldt dat het AI-systeem de laatste tijd een ongebruikelijk laag aantal ondeugdelijke producten detecteert. De reden voor de onderrapportage van ondeugdelijke producten is een software-update van het AI-systeem. Na handmatige inspectie blijken de producten ondeugdelijk te zijn en niet veilig te kunnen worden gebruikt.

In het rapport staat dat de valsnegatieven worden veroorzaakt door een cruciale fout in het nieuwe algoritme voor defectendetectie. De klokkenluider beweert dat managers op de hoogte waren van het probleem, maar het niet hebben aangepakt om reputatieschade voor het bedrijf te voorkomen.

Welke acties moeten nu worden ondernomen?

- A)** - Het interne algoritme aanpassen om het probleem aan te pakken  
- De relevante bevoegde autoriteit op de hoogte brengen als het probleem zich na 30 dagen nog steeds voordoet
- B)** - Het probleem intern onderzoeken en een start maken met het vinden van een oplossing  
- De relevante bevoegde autoriteit onmiddellijk op de hoogte brengen van het voorval
- C)** - De redenen onderzoeken waarom de klokkenluider de melding heeft gedaan  
- De relevante bevoegde autoriteit op de hoogte brengen als er klachten van consumenten binnenkomen
- D)** - Stoppen met het gebruik van het AI-systeem en overstappen op een oudere methode  
- Dan is het niet nodig om de relevante bevoegde autoriteit te informeren

### 15 / 40

MedTech Diagnostics gebruikt een AI-systeem met een hoog risico om medische diagnoses te stellen op basis van röntgenbeelden. Het bedrijf heeft de volgende voorzieningen getroffen:

- Het heeft met goed gevolg een externe audit doorlopen om te waarborgen dat het AI-systeem zich houdt aan de normen in de AI Act.
- Een robuust kader voor risicomanagement herkent en beperkt mogelijke problemen, waarvoor bovendien noodplannen beschikbaar zijn.
- Gedetailleerde gegevens over de activiteiten van het AI-systeem worden veilig opgeslagen voor verantwoordingsdoeleinden en audits.
- Gebruikers ontvangen duidelijke documentatie en trainingen waarin de besluitvorming en beperkingen van de AI worden uitgelegd.
- Alle door AI gegenereerde diagnoses worden beoordeeld door medische professionals voordat ze als definitief gelden, zodat menselijk toezicht is geïntegreerd.

Wat moet het bedrijf nog meer implementeren?

- A) Het moet robuuste procedures voor datagovernance toevoegen om de betrouwbaarheid en billijkheid van het AI-systeem te behouden.
- B) Het moet zorgen dat het AI-systeem onafhankelijk en zonder tussenkomst van een mens kan werken, zodat het de efficiëntie bevordert.
- C) Het moet een systeem implementeren om menselijke beslissingen automatisch te overrulen en zo het diagnostische proces te versnellen.
- D) Het moet een functie toevoegen die patiënten in staat stelt hun medisch dossier rechtstreeks te wijzigen op basis van AI-suggesties.

### 16 / 40

Een verzekeringsmaatschappij implementeert een nieuw kredietscoresysteem op basis van AI, met toegang tot zowel interne als openbare databases. De volgende risico's worden vastgesteld:

- **Een gebrek aan goede trainingsdata.** Als het model niet goed wordt getraind, is het lastig om mensen een nauwkeurige, eerlijke score te geven.
- **Integratie met andere toepassingen.** De AI-engine zal lastig te integreren zijn in de bestaande tamelijk complexe en op sommige punten verouderde toepassingsomgeving.
- **Niet-naleving van de AVG.** De Algemene Verordening Gegevensbescherming (AVG) bevat specifieke vereisten voor autonome verwerking van persoonsgegevens door geautomatiseerde systemen.
- **Transparantie en kwaliteit van het model.** Zowel medewerkers als klanten moeten de resultaten en beslissingen van het AI-model kunnen begrijpen.

De verzekeringsmaatschappij moet de AI Act naleven.

Welk risico is **niet** belangrijk voor naleving van de AI Act?

- A) Een gebrek aan goede trainingsdata
- B) Integratie met andere toepassingen
- C) Niet-naleving van de AVG
- D) Transparantie en kwaliteit van het model

**17 / 40**

Een overheidsinstantie stelt voor een AI-systeem te gebruiken om brandhaarden voor criminaliteit in het centrum van een grotere stad beter te kunnen voorspellen. Het systeem zal worden gebruikt voor geautomatiseerde surveillance. Het wordt zodanig geprogrammeerd dat iedereen die verdacht gedrag vertoont automatisch wordt geïdentificeerd en aan de lokale politie wordt doorgegeven. Dit is een geweldige kans om misdaad te voorkomen, het veiligheidsgevoel te vergroten en te zorgen dat niemand zijn straf ontloopt.

Kleven er risico's aan de implementatie van dit AI-systeem?

- A) Ja, want een risico op bias is inherent aan AI-systemen die worden gebruikt voor geautomatiseerde beslissingen, waardoor personen mogelijk onterecht worden benadeeld.
- B) Ja, want de AI Act voorziet zo veel privacyrisico's met bewakingssystemen dat de inzet ervan in openbare ruimtes ronduit verboden is onder deze wet.
- C) Nee, want voor de vervolging en preventie van misdaad vormen AI-systemen geen specifieke risico's, aangezien ze worden ingezet om de openbare veiligheid te verbeteren.
- D) Nee, want AI-systemen in het publieke domein vergroten de efficiëntie en vormen geen risico, aangezien alle beslissingen objectief en vrij van menselijke fouten zijn.

**18 / 40**

Een organisatie ontwikkelt een AI-systeem voor personeelswerving. Tijdens interne tests is het team op een risico gestuit: soms begunstigt het systeem onbedoeld sollicitanten met een bepaalde achtergrond, wat mogelijk tot discriminerende resultaten leidt. Het team twijfelt nu hoe het moet reageren op deze reden tot zorg.

De organisatie moet de AI Act naleven. Om het risico te beperken, hanteert de organisatie de norm ISO/IEC TR 24368.

Wat moet de organisatie volgens deze norm doen om dit risico te beperken?

- A) Het algoritme aanpassen zodat het prioriteit geeft aan demografische quota's op basis van werkgelegenheidsstatistieken
- B) Overstappen op een dataloze aanpak door alle demografische data uit de trainingsset te verwijderen
- C) Cyberbeveiligingsmaatregelen toepassen om de data van sollicitanten te beschermen en de integriteit van het systeem te verbeteren
- D) Een proces voor de betrokkenheid van stakeholders implementeren om mogelijke bias op te sporen en weg te nemen

**19 / 40**

Een organisatie ontwikkelt een AI-systeem om leningen goed te keuren. Tijdens interne tests stuit het nalevingsteam op deze risico's: de manier waarop het model beslissingen neemt is onvoldoende transparant en de documentatie voor het evalueren van risico's is beperkt.

De organisatie moet de AI Act naleven en hanteert daarvoor de norm ISO/IEC 42001 en het NIST AI Risk Management Framework (RMF).

Wat moet het bedrijf volgens deze norm en dit kader doen om deze risico's aan te pakken?

- A) Een beoordeling van de veiligheidsnaleving uitvoeren op basis van aanbevolen cyberbeveiligingsrichtlijnen
- B) Het AI-systeem onmiddellijk uit bedrijf nemen en overstappen op een handmatig goedkeuringsproces voor leningen
- C) Een meetplan met transparantiewaarden definiëren en voor toezichtdoeleinden de logica achter beslissingen vastleggen
- D) Het systeem opnieuw bouwen met synthetische data om zo veel mogelijk bronnen van bias weg te nemen

**20 / 40**

Een gerenommeerde autofabrikant heeft een verregaand geautomatiseerd voertuig (niveau 4) ontwikkeld, uitgerust met een op AI gebaseerde technologie voor objectherkenning om de verkeersveiligheid te waarborgen. Tijdens het testen worden de volgende risico's vastgesteld:

- Bij weinig licht is het systeem onvoldoende in staat om verkeersdrempels te herkennen.
- Mogelijk zal het model lastig te verkopen zijn, omdat het de afmetingen van andere voertuigen niet weet.
- De ontwikkelaars kunnen niet precies uitleggen hoe het model beslissingen neemt.
- In de ontwikkelingsfase van het AI-systeem zijn niet alle stakeholders betrokken bij het ontwikkelproces.

Wat moet hier worden aangepakt als er **alleen** naar naleving van de AI Act wordt gekeken?

- A) Het risico op onvoldoende tests in realistische omstandigheden
- B) Het risico op een gebrek aan transparantie in de besluitvorming door de AI
- C) Het risico op een beperkte schaalbaarheid van het AI-systeem naar andere voertuigmodellen
- D) Het risico op beperkte betrokkenheid van stakeholders tijdens de ontwikkeling van de AI

**21 / 40**

Een logistiek bedrijf werkt aan een AI-systeem om bezorgroutes te optimaliseren en brandstof te besparen. De organisatie overweegt twee mogelijkheden:

- Een closedsourcemodel van een leverancier die snellere installatie en geverifieerde nalevingscertificeringen garandeert.
- Een opensourcemodel dat meer aanpassingsmogelijkheden en transparantie biedt.

Het bedrijf moet de AI Act naleven, maar zoekt ook naar een goed evenwicht tussen innovatie en kosten.

Welk model past het **beste** bij dit bedrijf?

- A) Een closedsourcemodel, omdat zulke AI op zich veiliger is en meer wordt vertrouwd door instanties. Zo neemt de kans op niet-naleving af.
- B) Een closedsourcemodel, omdat dit vooraf gecertificeerde naleving biedt voor AI. Hierdoor hoeft het bedrijf minder moeite te doen om te bewijzen dat het de AI Act naleeft.
- C) Een opensourcemodel, omdat dit volledige transparantie voor de AI garandeert. Dat is gunstig voor de documentatie- en controleerbaarheidsvereisten.
- D) Een opensourcemodel, omdat dit is uitgesloten van naleving van de AI Act. Deze uitsluiting komt doordat de broncode openbaar beschikbaar is.

**22 / 40**

In de AI Act worden ethische principes voor AI-ontwikkeling gedefinieerd.

Wat is **geen** van deze principes?

- A) Verklaarbaarheid
- B) Billijkheid
- C) Verliespreventie
- D) Respect voor AI-autonomie

**23 / 40**

Een start-up ontwikkelt een AI-systeem om gepersonaliseerd leren op scholen te ondersteunen door lesplannen af te stemmen op de behoeften van individuele leerlingen. Het systeem verzamelt gegevens over de prestaties en het leergedrag van leerlingen.

Wat moet de start-up onder de AI Act overwegen om een evenwicht tussen innovatie en regelgeving te vinden?

- A) Het systeem niet als 'hoog risico' labelen om bijkomende regeldruk te omzeilen en innovatie te stroomlijnen
- B) Zorgen dat het AI-systeem een conformiteitsbeoordeling ondergaat en voldoet aan alle regelgeving voor systemen met een hoog risico
- C) Robuuste functies voor databescherming implementeren maar gebruikersmeldingen achterwege laten om vertragingen tijdens de inzet te vermijden
- D) Het systeem exclusief voor privéscholen in de handel brengen om de impact van nalevingsvereisten bij een hoog risico te beperken

**24 / 40**

De financiële instelling Fintegra implementeert een AI-systeem om frauduleuze transacties te detecteren. Voor analysedoeleinden moet het systeem toegang krijgen tot transactiegegevens en demografische gegevens van klanten. Fintegra moet voldoen aan het vereiste voor minimale gegevensverwerking uit de AI Act.

Hoe kan Fintegra het **beste** voldoen aan het vereiste voor minimale gegevensverwerking?

- A) De instelling moet alle transactiegegevens anonimiseren en alle data op basis waarvan de identiteit van natuurlijke personen kan worden herleid verwijderen, zelfs als die cruciaal is voor fraudedetectie.
- B) De instelling moet alle persoonsgegevens, inclusief volledige namen en exacte adressen, verzamelen om de precieze analyse en de verbetering na verloop van tijd te waarborgen en moet de data zo lang als nodig veilig bewaren.
- C) De instelling moet het verzamelen van gegevens beperken tot transactiegegevens die relevant zijn om fraude te detecteren, en vermijden dat persoonsgegevens zoals volledige namen of exacte adressen van klanten worden verwerkt.
- D) De instelling mag verzamelde gegevens alleen delen met erkende leveranciers die voldoen aan de AI Act, zodat de interne verwerking van persoonsgegevens zoals volledige namen van klanten tot een minimum beperkt blijft.

**25 / 40**

EduTech implementeert een adaptief leerplatform dat met behulp van AI leertrajecten voor leerlingen personaliseert. Het platform past het moeilijkheidsniveau van taken aan op basis van individuele prestaties.

Welk risico moet EduTech beperken om te waarborgen dat dit AI-systeem ethisch wordt gebruikt?

- A) Het risico op bias en discriminatie, want dit zou oneerlijke voor- of nadelen opleveren voor bepaalde leerlingen. Dit risico kan worden beperkt door datasets en algoritmes van het AI-systeem regelmatig te controleren en bij te werken.
- B) Het risico op een bovenmatige afhankelijkheid van technologie, waardoor leerlingen mogelijk geen kritische denkvaardigheden ontwikkelen. Dit risico kan worden beperkt door het besluitvormingsproces van het AI-systeem vertrouwelijk te houden, zodat leerlingen worden gestimuleerd om meer zelf na te denken.
- C) Het risico op inbreuken op de privacy, omdat gevoelige gegevens van leerlingen, waaronder hun prestaties, mogelijk verkeerd worden verwerkt of worden onthuld. Dit risico kan worden beperkt door meer aandacht te besteden aan betere technische prestaties van het AI-systeem.
- D) Het risico op transparantieproblemen, omdat leerlingen en onderwijzers mogelijk niet begrijpen hoe beslissingen worden genomen. Dit risico kan worden beperkt door te zorgen dat het AI-systeem werkt zonder menselijk toezicht, om zo de billijkheid te waarborgen.

**26 / 40**

Een ziekenhuisafdeling is gespecialiseerd in de diagnostiek en behandeling van ziektes. De afdeling ontwikkelt een diagnostisch AI-systeem dat moet helpen bij het opsporen van zeldzame aandoeningen. Het systeem analyseert patiëntgegevens, medische voorgeschiedenis en scans.

Het systeem wordt succesvol in gebruik genomen in de Verenigde Staten (VS). Enkele medisch specialisten in de Europese Unie (EU) willen het ook gaan gebruiken, maar begrijpen niet goed hoe het AI-systeem werkt. Bovendien beschikken ze niet over speciale kennis van en ervaring met AI-monitoring of het herkennen van storingen of foute diagnoses.

Wat is **geen** risico in verband met de overstap op dit AI-systeem?

- A) Het risico op ontbrekend effectief menselijk toezicht
- B) Het risico op onjuiste diagnoses wegens bias in de automatisering
- C) Het risico op wantrouwen wegens te weinig transparantie
- D) Het risico op onbevoegde inzage in patiëntendossiers

**27 / 40**

Een bedrijf maakt een AI-systeem voor slimme-woningassistenten. Tijdens het testen ontdekt het team dat fouten bij de spraakherkenning, zoals verwarrende woorden die ongeveer hetzelfde klinken, onbedoelde acties tot gevolg hebben, bijvoorbeeld dat het verkeerde apparaat wordt ingeschakeld. Deze fouten kunnen leiden tot inbreuken op de privacy, zoals gespreksopnames zonder toestemming of een onjuiste gebruikersidentificatie, waardoor gevoelige informatie mogelijk met onbevoegde partijen wordt gedeeld.

De organisatie moet de AI Act naleven en gebruikt daarvoor het kader CEN/CLC/TR 18115.

Wat moet de AI-aanbieder volgens dit kader doen om het beschreven probleem aan te pakken?

- A) Een plan voor de betrokkenheid van stakeholders opstellen om verschillende standpunten over de werking van het AI-systeem te verkrijgen
- B) Een ethische impactbeoordeling uitvoeren om inzicht te krijgen in de privacyrisico's van slimme-woningassistenten
- C) De kwaliteit van trainingsdata verbeteren door gebruik te maken van systematische validatie en methoden voor foutcontrole
- D) Door middel van versleuteling de dataveiligheid vergroten om spraakgegevens te beschermen en inbreuken te voorkomen

**28 / 40**

Een technologiebedrijf wordt betrapd op het gebruik van een AI-systeem dat in real time biometrische identificatie op afstand mogelijk maakt, wat nadrukkelijk verboden is onder de AI Act.

Wat is de juiste sanctie voor deze overtreding?

- A) Een formele waarschuwing zonder financiële sancties
- B) Een administratieve geldboete tot € 7,5 miljoen of 1% van de totale wereldwijde jaarlijkse omzet voor het voorafgaande boekjaar
- C) Een administratieve geldboete tot €15 miljoen of 3% van de totale wereldwijde jaarlijkse omzet voor het voorafgaande boekjaar
- D) Een administratieve geldboete tot €35 miljoen of 7% van de totale wereldwijde jaarlijkse omzet voor het voorafgaande boekjaar

29 / 40

De AI Act legt extra nadruk op het belang van twee aspecten van AI-systemen: transparantie en traceerbaarheid.

Waarom zijn transparantie en traceerbaarheid belangrijk?

- A) Omdat ze cruciaal zijn om verantwoording te waarborgen en het vertrouwen in AI-systemen te stimuleren
- B) Omdat dit verplichte vereisten zijn voor alle producten, inclusief AI-systemen
- C) Omdat ze van wezenlijk belang zijn voor de betrouwbaarheid en automatisering van AI-systemen
- D) Omdat de Europese, Chinese en Amerikaanse wetgeving deze aspecten gemeen hebben

30 / 40

Een retailorganisatie gebruikt een AI-systeem dat de manier waarop website-elementen worden weergegeven automatisch aanpast aan gebruikersvoorkeuren en het gebruikte apparaat. Het systeem geeft aanbevelingen voor producten en verbetert de gebruikerservaring op basis van de klikgeschiedenis en hoeveel tijd gebruikers op een pagina doorbrengen.

In welke categorie moet het gebruik van dit AI-systeem volgens de AI Act worden ingedeeld?

- A) Onaanvaardbaar risico
- B) Hoog risico
- C) Beperkt risico
- D) Laag of geen risico

31 / 40

Een bedrijf ontwikkelt een AI-systeem voor geautomatiseerde besluitvorming in het eigen wervingsproces. Het AI-systeem screent cv's, kent een score toe aan sollicitanten en geeft aanbevelingen voor sollicitatiegesprekken.

Het bedrijf maakt zich zorgen dat er in het AI-systeem sprake is van bias die mogelijk het wervingsproces beïnvloedt.

Wat kan het bedrijf het **beste** doen om bias in het AI-systeem te beperken?

- A) Het AI-systeem toestaan dat het leert en zich aanpast zonder verdere menselijke tussenkomst
- B) De bias in de trainingsdata negeren en focussen op de prestaties van het AI-systeem
- C) Een divers ontwikkelteam samenstellen om het AI-systeem te maken en monitoren
- D) Het AI-systeem trainen met één enkele databron om te zorgen voor consistentie

**32 / 40**

Feline Finesse is een webshop met accessoires en kussens voor kattenbezitters, die onder andere gepersonaliseerde knuffels op basis van ingestuurde foto's kunnen bestellen. De webshop gebruikt een AI-systeem dat het volgende kan:

- Het past prijzen dynamisch aan op basis van de activiteit van consumenten.
- Het rangschikt zoekresultaten op basis van klantvoorkeuren.
- Het geeft persoonlijke aanbevelingen voor andere producten waarvan het denkt dat klanten die leuk vinden.

De webshop wijst klanten momenteel op alles wat het AI-systeem doet, en is zeer transparant over hoe het algoritme werkt. De CEO zet echter vraagtekens bij deze praktijk en wil weten welke mate van transparantie vereist is en hoe die de omzet beïnvloedt.

Wat moet de CEO weten over transparantie binnen de context van de AI Act?

- A) Met transparantie kan de webshop aan consumenten bewijzen dat het systeem objectief is. Consumenten hebben onder de AI Act het recht om te begrijpen hoe hun data wordt gebruikt, en dit begrip bevordert hun vertrouwen.
- B) Transparantie kan de grenzen of beperkingen van het AI-systeem zichtbaar maken. Consumenten verliezen mogelijk hun vertrouwen in het bedrijf zodat zij dit beseffen en dit is schadelijk voor de reputatie van het bedrijf.
- C) Transparantie is niet verplicht voor e-commerce. Personalisering vergroot het gebruiksgemak voor consumenten en verder hoeven zij niet te weten of begrijpen hoe het AI-systeem werkt.
- D) Transparantie blijft beperkt tot het beschikbaar stellen van de broncode van het AI-systeem. Het consumentenvertrouwen in het systeem neemt mogelijk af als consumenten begrijpen hoe het algoritme precies werkt.

**33 / 40**

Voor een wervingsproces wordt een AI-systeem met een hoog risico gebruikt dat sollicitanten automatisch filtert op basis van hun kwalificaties. De gebruiksverantwoordelijke heeft echter geen mechanisme geïmplementeerd voor tussenkomst of toezicht door een mens in geval van twijfelachtige beslissingen.

Is menselijk toezicht vereist voor dit systeem onder de AI Act?

- A) Ja, want menselijk toezicht is nodig voor tussenkomst in besluitvormingsprocessen.
- B) Ja, want menselijk toezicht waarborgt dat verplichtingen op het gebied van billijkheid en transparantie worden nageleefd.
- C) Nee, want geautomatiseerde systemen zijn zodanig ontworpen dat ze functioneren zonder menselijke tussenkomst.
- D) Nee, want wervingsprocessen brengen geen kritieke veiligheidsrisico's voor natuurlijke personen met zich mee.

**34 / 40**

Een bedrijf treft voorbereidingen om een AI-model voor algemene doeleinden op de markt te brengen. Het model kan worden aangepast voor taken als automatisering van de klantenservice, contentcreatie en data-analyse. Het bedrijf is gevestigd buiten de Europese Unie (EU), maar is van plan het model in diverse EU-lidstaten te distribueren.

Wat is onder de AI Act **niet** vereist voordat het AI-model voor algemene doeleinden in de EU wordt gedistribueerd?

- A) Een gemachtigde in de EU aanstellen om alle nalevingskwesties af te handelen
- B) De EU-regelgeving op het gebied van auteursrechten naleven om het model te trainen met auteursrechtelijk beschermde data
- C) Een grondige audit uitvoeren om te verifiëren of alle EU-wetten en -regels volledig worden nageleefd
- D) Een gedetailleerde samenvatting publiceren van de gebruikte content om het AI-model voor algemene doeleinden te trainen

**35 / 40**

Een organisatie zet een AI-systeem in voor predictief onderhoud van industriële installaties. Wanneer het systeem enkele maanden in gebruik is, begint het een zeer hoog aantal valse meldingen te genereren, waardoor de workflows ontregeld raken. Uit een onderzoek blijkt het volgende:

- De organisatie heeft geen rekening gehouden met dynamische omgevingsveranderingen op de werkvloer.
- De organisatie beschikt niet over een formeel proces om risico's na de uitrol opnieuw te beoordelen.

De organisatie moet de AI Act naleven. Om deze kwesties op te lossen, hanteert de organisatie de norm ISO/IEC 23894.

Wat moet de organisatie volgens deze norm doen om de kwesties aan te pakken?

- A) Een workshop Human Centered Design (HCD) organiseren om de bruikbaarheid van het systeem te verbeteren
- B) Een proces voor risicomanagement met constante evaluatie en monitoring ontwikkelen
- C) Een audit van de cyberbeveiliging uitvoeren om mogelijke kwetsbaarheden op te sporen en aan te pakken
- D) Het AI-systeem vervangen door een eenvoudiger model op basis van regels dat makkelijker te beheersen is

**36 / 40**

Een wagenparkbeheerder gebruikt een AI-systeem om het gedrag van bestuurders bij te houden en voorspellingen over benodigd onderhoud te doen. Het systeem verzamelt en verwerkt grote datavolumes, zoals gps-locaties, rijpatronen en indicatoren voor voertuigprestaties. Tijdens een recente audit bleek dat het bedrijf beschikte over onvoldoende procedures voor databescherming.

Onder de AI Act zijn databeheer en privacybescherming essentieel voor dit bedrijf.

Waarom is dat zo?

- A) Omdat het bedrijf daarmee prioriteit kan geven aan zakelijke doelstellingen en operationele efficiëntie
- B) Omdat daarmee het gebruikersvertrouwen toeneemt, persoonsgegevens worden beschermd en onbevoegde inzage wordt voorkomen
- C) Omdat ze beide verplicht zijn en het bedrijf juridische problemen en mogelijke sancties kan voorkomen door de AI Act na te leven
- D) Omdat daarmee procedures voor dataverzameling worden gestroomlijnd, aangezien gebruikers niet langer toestemming hoeven te geven

**37 / 40**

Welk gebruik van een AI-systeem wordt onder de AI Act geclassificeerd als een beperkt risico?

- A) Een chatbot die is ontworpen om klanten te helpen met algemene vragen en die is geprogrammeerd om hen te laten weten dat het een AI-systeem is.
- B) Een systeem voor gezichtsherkenning dat wordt gebruikt om klanten in openbare ruimtes in real time te identificeren, bijvoorbeeld in een overdekt winkelcentrum.
- C) Een tool voor medische diagnoses om artsen te helpen om op basis van patiëntgegevens behandelingen aan te bevelen.
- D) Een AI-systeem dat een autonoom voertuig bedient dat zonder menselijk toezicht op de openbare weg rijdt.

**38 / 40**

Een bedrijf ontwikkelt een AI-systeem voor het onderwijs. Het AI-systeem bepaalt of een leerling toegang krijgt tot materialen, wordt toegelaten tot een school of wordt ingedeeld in een klas. Het AI-systeem zal beschikbaar zijn via clouddiensten.

In welke categorie moet het gebruik van dit AI-systeem volgens de AI Act worden ingedeeld?

- A) Onaanvaardbaar risico
- B) Hoog risico
- C) Beperkt risico
- D) Laag of geen risico

**39 / 40**

Een organisatie heeft een AI-systeem voor automatische personeelswerving ontwikkeld. Het team doet de volgende bevindingen tijdens het testen:

- Het systeem geeft sollicitanten met een bepaalde etnische achtergrond consistent een lagere score, omdat er sprake is van demografische bias in de trainingsdata.
- Er bestaat momenteel geen intern beoordelingsproces of feedbackmechanisme van relevante partijen dat op het risico op deze specifieke vorm van bias had kunnen wijzen.

De organisatie moet de AI Act naleven. Om deze kwesties op te lossen, hanteert de organisatie de norm ISO/IEC TR 24368.

Wat moet de organisatie volgens deze norm doen om de kwesties op te lossen?

- A) Een synthetische dataset maken om scheve demografische verhoudingen aan te pakken en de billijkheid te verbeteren
- B) Transparantie implementeren om de verklaarbaarheid en verantwoording van het systeem te vergroten
- C) De praktijken voor dataversleuteling aanscherpen en toegangsbeheer gebruiken om inbreuk te voorkomen
- D) Een ethisch kader met input van stakeholders gebruiken om mensenrechtenkwesties te evalueren

**40 / 40**

Een start-up ontwikkelt een AI-model voor algemene doeleinden dat wordt getraind met openbaar beschikbare onlinecontent, waaronder nieuwsberichten, onderzoeksrapporten en posts op social media. Na de lancering ontvangt het bedrijf een juridische kennisgeving van een groep auteurs die beweren dat hun intellectuele eigendommen zonder toestemming zijn gebruikt om het model te trainen.

Wat moet er in dit geval gedaan worden om intellectuele eigendomsrechten te beschermen?

- A) Aanvoeren dat het AI-model voor algemene doeleinden kan worden beschouwd als open source en daarom vrijgesteld is van nalevingsverplichtingen op het gebied van auteursrechten
- B) Beweren dat er sprake is van eerlijk gebruik onder de AI Act, aangezien de content openbaar beschikbaar was, en de dataset blijven gebruiken
- C) De door AI gegenereerde output die sterke overeenkomsten met de betwiste werken vertoont verwijderen om claims in verband van schendingen van het auteursrecht te vermijden
- D) Naleving waarborgen door details van de gebruikte trainingsdataset te documenteren en delen, inclusief de herkomst van de set

# Antwoordsleutel

1 / 40

De AI Act is een wet die is opgesteld voor de Europese Unie (EU). In artikel 1 worden de doelstellingen van de AI Act beschreven.

Wat zijn de **belangrijkste** doelstellingen van de AI Act?

- A) Richtlijnen die puur bedoeld zijn voor milieubescherming, zonder specifieke regels voor AI met een hoog risico, zonder verboden, en met innovatiemaatregelen alleen voor grote bedrijven
  - B) Geharmoniseerde regels voor AI-systemen in de EU, verboden voor bepaalde AI-praktijken, vereisten voor AI met een hoog risico, transparantieregels, markttoezicht en innovatieondersteuning
  - C) Verboden voor AI-praktijken, uitsluitend regels voor AI voor algemene doeleinden, transparantieregels behalve voor AI met een hoog risico en innovatieondersteuning beperkt tot niet-Europese entiteiten
  - D) Regels voor AI-systemen beperkt tot veiligheid en gezondheid, verboden voor alle AI-praktijken, transparantieregels uitsluitend voor AI met een hoog risico, en innovatieondersteuning voor alle bedrijven behalve start-ups
- 
- A) Incorrect. Bij deze optie wordt enkel gekeken naar milieubescherming en wordt voorbijgegaan aan specifieke regels voor AI met een hoog risico, verboden, en innovatiemaatregelen voor grote bedrijven. Dat is een verkeerde voorstelling van de brede benadering van de AI Act.
  - B) Correct. Deze optie geeft een juiste weerspiegeling van de hoofdpunten van de AI Act, waaronder geharmoniseerde regels voor AI-systemen, verboden voor bepaalde AI-praktijken, specifieke vereisten voor systemen met een hoog risico, transparantieregels, markttoezicht en innovatieondersteuning gericht op kleine en middelgrote ondernemingen (kmo's). (Literatuur: A, hoofdstuk 3.1, 3.2; AI Act, artikel 1)
  - C) Incorrect. Deze optie impliceert dat de AI Act alleen betrekking heeft op AI voor algemene doeleinden en sluit AI met een hoog risico uit van de transparantieregels, terwijl de innovatieondersteuning wordt beperkt tot niet-Europese entiteiten, wat niet conform de doelstellingen van de AI Act is.
  - D) Incorrect. De doelstellingen van de AI Act zijn breder dan alleen veiligheid en gezondheid. Bij deze optie wordt ten onrechte beweerd dat regels beperkt zijn tot veiligheid en gezondheid, worden start-ups uitgesloten van innovatieondersteuning en worden verboden voor alle AI-praktijken genoemd, wat niet in lijn is met de bepalingen van de AI Act.

2 / 40

Wat betekenen verantwoording en naleving volgens de AI Act?

- A) Verantwoording draait om het handhaven van gebruikersprivacy en dataveiligheid, terwijl naleving betrekking heeft op integratie met een bestaande IT-infrastructuur.
  - B) Verantwoording houdt in dat ontwikkelaars en operators die AI ontwikkelen daarvoor verantwoordelijk worden gehouden, en naleving betekent dat wettelijke vereisten in acht worden genomen.
  - C) Verantwoording draait om het garanderen dat AI-systemen winstgevend zijn voor ontwikkelaars, en naleving houdt in dat systemen voldoen aan de wensen en voorkeuren van gebruikers.
  - D) Verantwoording betekent dat AI-gebruikers verantwoordelijk worden gehouden voor correct gebruik van het systeem, terwijl naleving betekent dat binnen de sector gangbare normen voor AI-innovatie in acht worden genomen.
- 
- A) Incorrect. Hoewel gebruikersprivacy en dataveiligheid belangrijk zijn, zijn verantwoording en naleving onder de AI Act bredere concepten gericht op verantwoordelijkheid en de inachtneming van wettelijke en regelgevende vereisten, niet zozeer op enkel privacy- of integratiekwesties. Naleving draait om het opvolgen van wettelijke en ethische regels, niet om IT-integratie.
  - B) Correct. Verantwoording betekent dat ontwikkelaars en operators van AI-systemen verantwoordelijk kunnen worden gehouden voor hun handelingen en resultaten. Naleving verwijst naar het in acht nemen van wettelijke en regelgevende vereisten zoals beschreven in de AI Act om te waarborgen dat systemen veilig, transparant en billijk zijn. (Literatuur: A, hoofdstuk 3.10)
  - C) Incorrect. Verantwoording heeft niets te maken met winstgevendheid of gebruikersvoorkeuren, en naleving draait niet om het voldoen aan wensen van gebruikers. In plaats daarvan zijn beide gericht op de verantwoordelijkheid voor en inachtneming van wettelijke normen voor AI-systemen.
  - D) Incorrect. Verantwoording betekent onder andere dat ontwikkelaars en operators verantwoordelijk worden gehouden voor de acties van AI-systemen, niet dat de schuld wordt afgeschoven. Naleving draait meer om het voldoen aan specifieke wettelijke en regelgevende normen dan om algemene normen binnen een sector.

3 / 40

Personen die de gevolgen ondervinden van AI-systemen hebben onder de AI Act specifieke rechten om transparantie, billijkheid en verantwoording te waarborgen.

Welk recht wordt expliciet verleend onder de AI Act?

- A) Het recht om te worden geïnformeerd over het feit dat ze interageren met of de gevolgen ondervinden van een AI-systeem
  - B) Het recht om inzage in de broncode van het AI-systeem te eisen
  - C) Het recht om het gebruik van AI te verbieden voor elk besluitvormingsproces dat betrekking op hen heeft
  - D) Het recht om verwijdering van door het AI-systeem gebruikte persoonsgegevens aan te vragen
- 
- A) Correct. Personen moeten ervan in kennis worden gesteld dat zij interageren met een AI-systeem, tenzij dit duidelijk is voor een redelijk opmerkzaam persoon. Zo wordt transparantie gewaarborgd en begrijpen personen wanneer AI van invloed is op beslissingen waardoor zij mogelijk de gevolgen van ondervinden. (Literatuur: A, hoofdstuk 3.6; AI Act, artikel 50)
  - B) Incorrect. De AI Act geeft personen geen recht op inzage in de broncode van een AI-systeem. De transparantieverplichtingen zijn gericht op uitleg en bekendmaking, niet op volledige inzage in bedrijfseigen code.
  - C) Incorrect. De AI Act geeft personen niet het recht om de inzet van AI bij besluitvorming te verbieden, maar zorgt wel dat daarbij sprake is van toezicht en transparantie.
  - D) Incorrect. Hoewel databeschermingswetten personen bepaalde rechten met betrekking tot hun data bieden, verleent de AI Act geen algemeen recht om te vragen om verwijdering van alle data die door een AI-systeem wordt gebruikt.

4 / 40

Anna is als compliance officer bij een kleine of middelgrote onderneming (kmo) verantwoordelijk voor het toezicht op de implementatie van een nieuw AI-systeem dat zal worden ingezet om de klantondersteuning te automatiseren. Het bedrijf heeft dit systeem niet zelf gebouwd, maar koopt het van een andere aanbieder. Onder de AI Act wordt dit systeem geclassificeerd als een AI-systeem met een hoog risico.

Anna is gevraagd om te zorgen dat het bedrijf bij de inzet en monitoring van dit AI-systeem voldoet aan de verplichtingen van de gebruiker. Ze moet bepalen welke acties prioriteit moeten krijgen en welke acties moeten worden vermeden.

Wat kan Anna beter **niet** overwegen, gezien de verplichtingen voor AI-gebruikers?

- A) De algoritmes van het AI-model verder ontwikkelen om de besluitvormende capaciteiten te verbeteren zonder de aanbieder daarbij te betrekken
  - B) De prestaties van het AI-systeem gedetailleerd bijhouden en zorgen dat aan alle relevante rapportagevereisten wordt voldaan
  - C) De prestaties van het AI-systeem monitoren om te waarborgen dat het werkt zoals beoogd en voldoet aan veiligheidsnormen
  - D) Elk ernstig incident met of gebrekkig functioneren van het AI-systeem melden bij de toepasselijke instanties, zoals wettelijk vereist
- 
- A) Correct. Anna mag niet proberen het algoritme van het AI-systeem onafhankelijk te ontwikkelen of de besluitvormende capaciteiten ervan te wijzigen zonder de aanbieder daarbij te betrekken. Dit valt buiten de scope van de verplichtingen van de gebruiker en kan leiden tot niet-naleving of andere onbedoelde gevolgen. Gebruikers zijn niet verantwoordelijk voor aanpassingen aan de interne structuur van het systeem. (Literatuur: A, hoofdstuk 1.1)
  - B) Incorrect. Gegevens bijhouden en de transparantie waarborgen zijn onder de AI Act in lijn met de wettelijke vereisten voor gebruikers van systemen met een hoog risico. Dit komt de verantwoording en regelnaleving ten goede.
  - C) Incorrect. Onder de AI Act is monitoring van prestaties een belangrijke verplichting van gebruikers om te waarborgen dat de AI werkt zoals beoogd en geen veiligheidsrisico's vormt.
  - D) Incorrect. Melden van gebrekkig functioneren of ernstige incidenten is onder de AI Act een belangrijke vereiste voor gebruikers van systemen met een hoog risico, omdat zo de naleving kan worden gehandhaafd en potentiële risico's snel kunnen worden aangepakt.

5 / 40

Een AI-systeem voor gezichtsherkenning wordt in openbare ruimtes gebruikt voor beveiligingsdoeleinden. Eén organisatie is voor dit AI-systeem het meest relevant voor toezicht op de naleving van alle regelgeving op het gebied van databescherming en privacy (zoals de Algemene Verordening Gegevensbescherming (AVG)).

Welke organisatie is dat?

- A) De Europese Consumentenorganisatie (Bureau Européen des Unions de Consommateurs, BEUC)
  - B) De Europese Raad voor Artificiële Intelligentie (European Artificial Intelligence Board; AI-board)
  - C) Het Europese Hof van Justitie (EHJ)
  - D) Het Europese Comité voor Gegevensbescherming (EDPB)
- 
- A) Incorrect. De BEUC richt zich op consumentenrechten, die raakvlakken hebben met biometrie en databescherming voor AI-specifieke regelgeving.
  - B) Incorrect. De AI-board houdt toezicht op naleving van de AI Act en richt zich daarbij op AI-specifieke regelgeving. Deze vraag heeft echter betrekking op databescherming en privacy, en die vallen onder de AVG.
  - C) Incorrect. Het Europese Hof van Justitie behandelt juridische geschillen, maar houdt geen toezicht op AI-regelgeving of de verwerking van biometrische gegevens.
  - D) Correct. Het EDPB is verantwoordelijk voor een consistente toepassing van de AVG in alle lidstaten van de Europese Unie (EU). Samen met nationale autoriteiten voor databescherming pakt het kwesties aan zoals het gebruik van biometrische data in AI-beveiligingssystemen. (Literatuur: A, hoofdstuk 3.9, 3.10, 4.5)

**6 / 40**

Een bedrijf ontwikkelt een AI-systeem voor gepersonaliseerde marketing. Dit systeem maakt gebruik van machine learning algoritmes (ML-algoritmes) om advertenties af te stemmen op individuele klanten. Tijdens een nalevingsbeoordeling stelt het team de volgende risico's vast:

- Er is geen documentatie waaruit duidelijk blijkt hoe het AI-systeem data verwerkt.
- Het proces aan de hand waarvan het AI-systeem gepersonaliseerde aanbevelingen doet, wordt niet volledig begrepen.
- Klanten klagen over deze kwesties.

Het bedrijf moet de AI Act naleven en hanteert daarvoor de norm ISO/IEC 42001 en het NIST AI Risk Management Framework (RMF).

Wat moet het bedrijf volgens deze norm en dit kader doen om deze kwesties aan te pakken?

- A) Een reeks tests van de gebruikerservaring (UX) uitvoeren om feedback over de bruikbaarheid, leerbaarheid en klantvoorkeuren te verzamelen
  - B) Proberen om voorspellingen van het systeem nauwkeuriger te maken om zo de kostenefficiëntie, klanttevredenheid en betrokkenheid te verbeteren
  - C) Een documentatieproces implementeren waarin databronnen, verwerkingsmethoden en de algoritmische besluitvorming uitvoerig worden beschreven
  - D) De systeemhardware upgraden om snellere verwerking, meer efficiëntie en een hogere klanttevredenheid te waarborgen
- 
- A) Incorrect. Tests van UX bieden waardevolle inzichten in de effectiviteit van het systeem, maar daarmee is de onderliggende kwestie van ontbrekende documentatie en een gebrekkige transparantie binnen zowel het data- als het besluitvormingsproces niet opgelost.
  - B) Incorrect. Hoewel nauwkeurigere voorspellingen de tevredenheid van consumenten kunnen verbeteren, worden daarmee de hoofdkwesties van de documentatie en transparantie bij de dataverwerking en besluitvorming niet aangepakt.
  - C) Correct. Implementatie van een uitgebreid documentatieproces is in lijn met de norm ISO/IEC 42001, waarin transparantie en documentatie voor de gehele AI-levenscyclus wordt benadrukt. Het NIST AI RMF ondersteunt dit ook door het gedetailleerd bijhouden van data en beslissingen te stimuleren om zo traceerbaarheid en verantwoording te waarborgen. (Literatuur: B, hoofdstuk 2.3)
  - D) Incorrect. Met een hardware-upgrade neemt de verwerkingssnelheid misschien wel toe, maar worden de kwesties op het gebied van transparantie en benodigde documentatie voor naleving van de AI Act niet aangepakt.

7 / 40

Een bedrijf ontwikkelt een AI-systeem voor de monitoring van patiënten die in een ziekenhuis worden opgenomen. Het systeem werkt met high-definition camera's (HD-camera's) op de kamers om de status van patiënten in real time te monitoren. Als het systeem detecteert dat een patiënt onrustig is, wordt er automatisch iemand van de verpleging opgeroepen om langs te gaan.

Om de prestaties van het AI-systeem te verbeteren, wil het bedrijf een videodatabase van patiënten aanleggen, waarbij video's op kritieke punten worden voorzien van een opmerking van een professional. Zo komt er meer trainingsdata voor het systeem beschikbaar.

Het bedrijf overweegt een gegevensbeschermingseffectbeoordeling (DPIA) uit te voeren. Het verantwoordelijke team twijfelt of een DPIA nodig is. Als een DPIA verplicht is, wil het team weten wanneer de beoordeling dient te worden uitgevoerd: nu of pas nadat de update is uitgerold.

Het bedrijf moet de AI Act en de Algemene Verordening Gegevensbescherming (AVG) naleven.

Moet het bedrijf nu al een DPIA uitvoeren?

- A) Ja, want een DPIA is verplicht voor AI-projecten die mogelijk een hoog risico vormen voor de rechten van natuurlijke personen.
  - B) Ja, want een DPIA is verplicht voor elk project waarbij persoonsgegevens worden verzameld, zelfs als het risico van het project laag is.
  - C) Nee, want een DPIA is niet verplicht wanneer gegevens worden gebruikt voor trainingsdoeleinden, onderwijs of wetenschappelijk onderzoek.
  - D) Nee, want een DPIA is pas verplicht nadat een AI-systeem volledig is ontwikkeld, getest en uitgerold.
- 
- A) Correct. Deze optie weerspiegelt de vereisten onder de AVG. Een DPIA is verplicht wanneer verwerkingsactiviteiten waarschijnlijk resulteren in een hoog risico voor de rechten en vrijheden van natuurlijke personen, vooral bij gebruik van nieuwe technologieën zoals AI-systemen die (zeer) gevoelige gegevens zoals patiëntvideo's verwerken. Deze vallen onder de categorie gezondheidsgegevens, waarvoor extra waarborgen vereist zijn (Literatuur: A, hoofdstuk 4.5)
  - B) Incorrect. Hoewel een DPIA belangrijk is, is deze niet automatisch verplicht voor elk project waarbij persoonsgegevens worden verzameld. De AVG schrijft een DPIA vooral voor wanneer de gegevensverwerking waarschijnlijk resulteert in hoge risico's voor de rechten en vrijheden van personen, niet simpelweg wegens het verzamelen van persoonsgegevens, zoals biometrische gegevens, gezondheidsgegevens of monitoring op grote schaal.
  - C) Incorrect. Het doel van het datagebruik (bijvoorbeeld training of onderzoek) kan een project niet vrijwaren van de verplichting om een DPIA uit te voeren. De AVG is nog steeds van toepassing, vooral wanneer de verwerking mogelijk resulteert in hoge risico's voor personen en met name wanneer het gevoelige gegevens zoals patiëntvideo's betreft.
  - D) Incorrect. Om risico's proactief op te sporen en te beperken moet een DPIA worden uitgevoerd voorafgaand aan de verwerking, met name tijdens de plannings- en ontwikkelingsfase van een project. Wanneer er wordt gewacht tot na de uitrol, leidt dit mogelijk tot niet-naleving van de AVG.

8 / 40

Een bedrijf ontwikkelt een AI-systeem voor gezichtsherkenning in real time. Een particulier beveiligingsbedrijf zet het AI-systeem in voor de monitoring van een openbaar overdekt winkelcentrum. Het systeem scant alle bezoekers, vergelijkt hen met een database met mensen die eerder een overtreding hebben begaan en een database met politiek activisten, en markeert bezoekers die voorkomen in een van beide databases. Gemarkeerde bezoekers worden gedurende hun hele bezoek onopgemerkt gevolgd om te beoordelen of zij zich volgens het beveiligingsbedrijf verdacht gedragen.

In welke categorie moet het gebruik van dit AI-systeem volgens de AI Act worden ingedeeld?

- A) Onaanvaardbaar risico
  - B) Hoog risico
  - C) Beperkt risico
  - D) Laag of geen risico
- 
- A) Correct. Openbare biometrische identificatie in real time om personen te volgen op basis van hun politieke activiteit is verboden onder artikel 5 van de AI Act. Onder de wet is het niet toegestaan AI-systemen te gebruiken voor niet-gerichte surveillance en beoordeling van natuurlijke personen (social scoring) als dit in strijd is met grondrechten. (Literatuur: A, hoofdstuk 3.3, 3.4; AI Act, artikel 5(1)(d))
  - B) Incorrect. Het systeem valt niet slechts in de categorie hoog risico, maar in de categorie onaanvaardbaar risico, omdat het personen volgt op basis van hun politieke activiteit en zonder hun toestemming. Dit is een verboden manier om een AI-systeem te gebruiken.
  - C) Incorrect. Onder een beperkt risico vallen systemen als chatbots of instrumenten voor het bepalen van de gemoedstoestand met transparantieplichtingen. Omdat er in dit geval sprake is van biometrische surveillance met ernstige implicaties voor de rechten van personen, is het risico niet beperkt.
  - D) Incorrect. Een laag risico is van toepassing op systemen met weinig impact, zoals spamfilters. Gezichtsherkenning om mensen te volgen overtreft dit risico en is nadrukkelijk verboden.

9 / 40

Een reisbureau gebruikt een AI-systeem om dynamische, gerichte marketingcampagnes voor pakketreizen te ontwikkelen. De campagnes omvatten realtimeadvertenties op social media en reisplatforms, waarbij gebruik wordt gemaakt van de browsegeschiedenis van sitebezoekers. Het reisbureau gebruikt AI om de gemoedstoestand van gebruikers te bepalen en stemt vervolgens de suggesties voor bestemmingen en activiteiten hierop af.

Het reisbureau moet de AI Act naleven.

Welk risico moet het reisbureau aanpakken?

- A) Het risico op mogelijke bias. Het bureau moet de trainingsdata regelmatig bijwerken om te voorkomen dat er suggesties voor irrelevante bestemmingen verschijnen of een onjuiste gemoedstoestand wordt bepaald.
  - B) Het risico op ineffectieve reclameactiviteiten. Het bureau moet vooral proberen het algoritme bij te werken, want gepersonaliseerde advertenties vallen niet onder de AI Act.
  - C) Het risico op een gebrek aan transparantie. Het bureau moet openheid over de AI garanderen, bias in suggesties terugdringen en beoordelen of de reclameactiviteiten ethisch zijn.
  - D) Het risico op misbruik van persoonsgegevens. Het bureau moet stoppen met het gebruik van AI-gestuurde personalisatie omdat het onder de AI Act verboden is om persoonsgegevens te gebruiken voor gerichte advertenties.
- 
- A) Incorrect. Bias onder de AI Act is bias waarbij bepaalde klantgroepen worden benadeeld. Het is onwaarschijnlijk dat een irrelevante suggestie voor een reisbestemming veel impact heeft op klanten.
  - B) Incorrect. Hoewel de AI Act AI-toepassingen met een hoog risico de hoogste prioriteit geeft, is de wet ook van toepassing op commerciële sectoren zoals reclame en toerisme, vooral wanneer er sprake is van klantprofilering en besluitvorming.
  - C) Correct. Dit zijn de voornaamste verplichtingen onder de AI Act. Bureaus moeten transparant zijn door consumenten te informeren over de inzet van AI, bias te voorkomen en te garanderen dat reclamestrategieën aan ethische normen voldoen. (Literatuur: A, hoofdstuk 7.8, 7.9)
  - D) Incorrect. Onder de AI Act zijn zowel advertenties op maat als AI-gestuurde personalisatie niet nadrukkelijk verboden. De wet bevat eerder richtlijnen voor moreel gedrag waarin wordt opgeroepen tot openheid, rechtvaardigheid en dataveiligheid om te zorgen dat gebruikte methoden voldoen aan de waarden van de Europese Unie (EU).

10 / 40

Een bedrijf heeft een AI-model ontwikkeld dat in verschillende sectoren kan worden gebruikt, waaronder de gezondheidszorg en financiële dienstverlening. Vanwege de brede toepassingsmogelijkheden brengt het AI-model mogelijk risico's voor de volksgezondheid met zich mee.

Wat heeft het bedrijf ontwikkeld en welke praktijken dient het bedrijf volgens de AI Act te implementeren?

- A) Het bedrijf heeft een AI-systeem voor algemene doeleinden ontwikkeld dat systeemrisico's met zich meebrengt. Het dient aanvullende tests uit te voeren om die risico's te beperken.
  - B) Het bedrijf heeft een AI-systeem met een hoog risico ontwikkeld. Het dient alle vereisten voor AI-systemen met een hoog risico te implementeren zoals beschreven in de AI Act.
  - C) Het bedrijf heeft een beperkt (narrow) AI-model ontwikkeld. Om risico's te voorkomen dient het te waarborgen dat het model alleen werkt binnen vooraf gedefinieerde parameters.
  - D) Het bedrijf heeft een experimenteel AI-model ontwikkeld. Het dient zich op onderzoek en ontwikkeling te richten zonder onmiddellijk risicomanagement.
- 
- A) Correct. Het bedrijf heeft een AI-systeem voor algemene doeleinden ontwikkeld dat systeemrisico's met zich meebrengt. Het bedrijf dient een aanvullende beoordeling van het model uit te voeren aan de hand van protocollen en geavanceerde instrumenten, waaronder de implementatie en documentatie van beveiligingstests. (Literatuur: A, hoofdstuk 3.7; AI Act, artikel 3 en 55)
  - B) Incorrect. Hoewel het AI-model mogelijke risico's met zich meebrengt, wordt het geclassificeerd als een AI-systeem voor algemene doeleinden, niet specifiek als een systeem met een hoog risico.
  - C) Incorrect. Een beperkt AI-model voert slechts enkele specifieke taken uit waaraan niet dezelfde systeemrisico's kleven als aan een AI-systeem voor algemene doeleinden.
  - D) Incorrect. Zelfs experimentele AI-modellen dienen zich aan praktijken voor risicomanagement te houden als ze mogelijke systeemrisico's vormen.

**11 / 40**

Een organisatie ontwikkelt een AI-systeem met een hoog risico. Tijdens het testen stelt het ontwikkelteam diverse risico's vast, waaronder inconsistenties in de volledigheid van data en de aanwezigheid van verouderde records. Deze risico's hebben mogelijk een negatieve impact op de prestaties van het model.

De organisatie moet de AI Act naleven en gebruikt daarvoor het kader CEN/CLC/TR 18115.

Wat moet de organisatie volgens dit kader doen om deze risico's aan te pakken?

- A)** Een gegevensbeschermingseffectbeoordeling (DPIA) uitvoeren om de billijkheid van besluitvorming door AI aan te pakken
  - B)** Alle datasets voor trainings- en testdoeleinden versleutelen aan de hand van protocollen om onbevoegde inzage in persoonsgegevens te voorkomen
  - C)** Algemene maatregelen voor risicobeheer implementeren om de genoemde operationele risico's en reputatierisico's te beperken
  - D)** De datakwaliteit verbeteren door gestructureerde meetwaarden voor de kwaliteit en statistische evaluatiemethoden toe te passen
- 
- A)** Incorrect. Een DPIA is nuttig om risico's voor de rechten en vrijheden van personen te beoordelen, maar is niet het juiste instrument om problemen met verouderde of onvolledige data rechtstreeks aan te pakken. Een betere datakwaliteit vraagt om technische maatregelen, geen juridische beoordeling.
  - B)** Incorrect. Hoewel versleuteling belangrijk is voor de dataveiligheid, worden daarmee geen problemen met de datakwaliteit aangepakt, zoals volledigheid en tijdigheid. Artikel 10 van de AI Act benadrukt niet alleen de bescherming van data, maar waarborgt ook dat de data die wordt gebruikt om AI te trainen en testen relevant, representatief en van goede kwaliteit is.
  - C)** Incorrect. Hoewel risicomanagement zeer belangrijk is bij AI, wordt daarmee niet het benodigde kader voor datakwaliteit geboden om de vastgestelde problemen op te lossen. De volledigheid en tijdigheid van data kan worden beheerd door de datakwaliteit te verbeteren, niet met algemene normen voor AI-risico's.
  - D)** Correct. CEN/CLC/TR 18115 bevat richtlijnen om de datakwaliteit in de hele levenscyclus van AI-systemen te evalueren en verbeteren. De nadruk ligt daarbij op het gebruik van meetwaarden voor kenmerken als volledigheid en tijdigheid, vooral tijdens het voorbereiden van data, om zo naleving van artikel 10 van de AI Act te waarborgen. (Literatuur: B, hoofdstuk 1; AI Act, artikel 10)

12 / 40

In de AI Act worden diverse rollen gerelateerd aan een AI-systeem beschreven.

Wat is de definitie van de rol 'importeur van een AI-systeem'?

- A) Een persoon of organisatie die een AI-systeem ontwerpt, ontwikkelt en onder diens eigen naam of merk in de handel brengt.
  - B) Een persoon of organisatie die een AI-systeem in de handel brengt, maar niet verantwoordelijk is voor de oorspronkelijke ontwikkeling ervan.
  - C) Een persoon of organisatie die een AI-systeem gebruikt voor eigen activiteiten en zorgt dat alle lokale verplichtingen van de gebruiker worden nageleefd.
  - D) Een regelgevende instantie die als taak heeft te monitoren of het AI-systeem wordt geïmporteerd conform de bepalingen van de AI Act.
- 
- A) Incorrect. Dit is een omschrijving van de rol 'aanbieder', die verantwoordelijk is voor de ontwikkeling en in de handel brengen van het AI-systeem, niet voor de import ervan.
  - B) Correct. Dit is de definitie van de rol 'importeur', doorgaans wanneer het systeem buiten de Europese Unie (EU) is ontwikkeld. Het is de verantwoordelijkheid van de importeur om te zorgen dat het systeem voldoet aan alle regelgevende vereisten van de EU en om samen met de aanbieder deze naleving aan te tonen. (Literatuur: A, hoofdstuk 3.1)
  - C) Incorrect. Dit is een beschrijving van de rol 'gebruiker', degene die werkt met het AI-systeem en het monitort, niet van de importeur.
  - D) Incorrect. Regelgevende instanties zijn geen importeurs. Het is hun verantwoordelijkheid om naleving van de wet af te dwingen, maar ze zijn niet actief betrokken bij het in de handel brengen van systemen.

**13 / 40**

Een bedrijf ontwikkelt een AI-systeem voor fraudedetectie bij financiële transacties. Dit systeem analyseert transactiepatronen om verdachte activiteiten op te sporen en frauduleus gedrag te voorkomen. Gezien de kans op valspositieven die mogelijk een nadelige invloed op legitieme transacties hebben en het feit dat fraudetactieken voortdurend veranderen, erkent het bedrijf dat effectieve waarborgen vereist zijn.

Het bedrijf moet de AI Act naleven. Om problemen met valspositieven te helpen voorkomen, hanteert het bedrijf de norm ISO/IEC 23894.

Wat moet het bedrijf volgens deze norm doen om deze problemen te voorkomen?

- A)** Risicomanagement integreren in alle activiteiten om breed toezicht en proactieve risicobeperking te waarborgen
  - B)** Maatregelen voor dataprivacy uitbreiden om gevoelige informatie te beschermen en te voldoen aan privacyregels
  - C)** Proberen de nauwkeurigheid van het model te verbeteren om betrouwbare prestaties te waarborgen en valspositieven tot een minimum te beperken
  - D)** Cyberbeveiligingsmaatregelen implementeren om het systeem te beschermen tegen dreigingen van buitenaf en onbevoegde inzage
- 
- A)** Correct. Bij deze aanpak wordt risicomanagement geïntegreerd in de gehele organisatie door kaders aan te passen aan specifieke contexten en stakeholders te betrekken, zodat AI-gerelateerde risico's effectief worden opgespoord en beperkt. Dit is in lijn met de norm ISO/IEC 23894 en helpt het beste om naleving te bewerkstelligen. (Literatuur: B, hoofdstuk 3.2)
  - B)** Incorrect. Het bedrijf moet belangrijke privacykwesties aanpakken. Daarmee worden echter niet specifiek de benodigde uitgebreide risicomanagementpraktijken voor AI geïntegreerd zoals geschetst in norm ISO/IEC 23894, wat wel nodig is voor naleving.
  - C)** Incorrect. Door te proberen de nauwkeurigheid van het model te verbeteren om betrouwbare prestaties te waarborgen en valspositieven tot een minimum te beperken, verbetert het bedrijf wel de effectiviteit van het model, maar pakt het niet de bredere aspecten van risicomanagement aan, zoals het opsporen, beoordelen en beperken van potentiële AI-gerelateerde risico's. De norm ISO/IEC 23894 benadrukt een uitgebreide aanpak van risicomanagement.
  - D)** Incorrect. Door cyberbeveiligingsmaatregelen te implementeren om het systeem te beschermen tegen dreigingen van buitenaf en onbevoegde inzage, pakt het bedrijf een belangrijk onderdeel van de systeembeveiliging aan. Dit omvat echter niet de volledige scope van risicomanagementpraktijken voor AI, zoals gespecificeerd in de norm ISO/IEC 23894.

14 / 40

Een productiebedrijf maakt in zijn assemblagelijnen gebruik van AI-gestuurde robotsystemen voor kwaliteitscontrole. Het onderzoeksteam verneemt dat een anonieme klokkenluider meldt dat het AI-systeem de laatste tijd een ongebruikelijk laag aantal ondeugdelijke producten detecteert. De reden voor de onderrapportage van ondeugdelijke producten is een software-update van het AI-systeem. Na handmatige inspectie blijken de producten ondeugdelijk te zijn en niet veilig te kunnen worden gebruikt.

In het rapport staat dat de valsnegatieven worden veroorzaakt door een cruciale fout in het nieuwe algoritme voor defectendetectie. De klokkenluider beweert dat managers op de hoogte waren van het probleem, maar het niet hebben aangepakt om reputatieschade voor het bedrijf te voorkomen.

Welke acties moeten nu worden ondernomen?

- A) - Het interne algoritme aanpassen om het probleem aan te pakken  
- De relevante bevoegde autoriteit op de hoogte brengen als het probleem zich na 30 dagen nog steeds voordoet
  - B) - Het probleem intern onderzoeken en een start maken met het vinden van een oplossing  
- De relevante bevoegde autoriteit onmiddellijk op de hoogte brengen van het voorval
  - C) - De redenen onderzoeken waarom de klokkenluider de melding heeft gedaan  
- De relevante bevoegde autoriteit op de hoogte brengen als er klachten van consumenten binnenkomen
  - D) - Stoppen met het gebruik van het AI-systeem en overstappen op een oudere methode  
- Dan is het niet nodig om de relevante bevoegde autoriteit te informeren
- 
- A) Incorrect. Als het probleem wordt aangepakt zonder een autoriteit in te lichten, wordt er niet voldaan aan de wettelijke verplichting om ernstige incidenten te melden. Niet voldoen aan de meldingsplicht kan leiden tot sancties en wantrouwen over de manier waarop het bedrijf AI-systemen inzet.
  - B) Correct. Door onderzoek te doen wordt de onderliggende oorzaak van het probleem gevonden en aangepakt. Met een melding aan de relevante bevoegde autoriteit wordt naleving van de wet gegarandeerd. (Literatuur: A, hoofdstuk 7.4, 3.10; AI Act, artikel 73)
  - C) Incorrect. Onderzoek naar de motieven van een klokkenluider is een schending van de beschermingsprincipes voor klokkenluiders en heeft een ontmoedigend effect op ethische meldingen. Bij deze aanpak krijgt reputatiebeheer prioriteit boven wettelijke en ethische verplichtingen, wat een overtreding is van de vereisten in de AI Act en de integriteit van de organisatie schaadt.
  - D) Incorrect. Het probleem is weliswaar verholpen wanneer het AI-systeem wordt gestopt, maar daarmee wordt nog niet voldaan aan de meldingscriteria uit de AI Act. Deze keuze is onacceptabel, omdat dit een ernstig incident met gevolgen voor de productveiligheid was, dat moet worden gemeld en aangepakt.

## 15 / 40

MedTech Diagnostics gebruikt een AI-systeem met een hoog risico om medische diagnoses te stellen op basis van röntgenbeelden. Het bedrijf heeft de volgende voorzieningen getroffen:

- Het heeft met goed gevolg een externe audit doorlopen om te waarborgen dat het AI-systeem zich houdt aan de normen in de AI Act.
- Een robuust kader voor risicomanagement herkent en beperkt mogelijke problemen, waarvoor bovendien noodplannen beschikbaar zijn.
- Gedetailleerde gegevens over de activiteiten van het AI-systeem worden veilig opgeslagen voor verantwoordingsdoeleinden en audits.
- Gebruikers ontvangen duidelijke documentatie en trainingen waarin de besluitvorming en beperkingen van de AI worden uitgelegd.
- Alle door AI gegenereerde diagnoses worden beoordeeld door medische professionals voordat ze als definitief gelden, zodat menselijk toezicht is geïntegreerd.

Wat moet het bedrijf nog meer implementeren?

- A) Het moet robuuste procedures voor datagovernance toevoegen om de betrouwbaarheid en billijkheid van het AI-systeem te behouden.
  - B) Het moet zorgen dat het AI-systeem onafhankelijk en zonder tussenkomst van een mens kan werken, zodat het de efficiëntie bevordert.
  - C) Het moet een systeem implementeren om menselijke beslissingen automatisch te overrulen en zo het diagnostische proces te versnellen.
  - D) Het moet een functie toevoegen die patiënten in staat stelt hun medisch dossier rechtstreeks te wijzigen op basis van AI-suggesties.
- 
- A) Correct. Met robuuste data en datagovernance zorgt het bedrijf voor kwaliteit, beperkt het bias en bevordert het de traceerbaarheid van trainings- en gebruiksdata, wat het systeem eerlijk en betrouwbaar maakt. (Literatuur: A, hoofdstuk 3.3; AI Act, artikel 15)
  - B) Incorrect. Volledige automatisering bij het stellen van medische diagnoses is niet toegestaan onder de AI Act. Voor AI-systemen met een hoog risico in de gezondheidszorg is menselijk toezicht vereist om veiligheid en nauwkeurigheid te waarborgen, wat volledige onafhankelijkheid ongepast maakt. Beoordeling door een mens is essentieel met het oog op de patiëntveiligheid en regelnaleving.
  - C) Incorrect. Als menselijke beslissingen automatisch worden overruled, komt de patiëntveiligheid in het gedrang en wordt de onmisbare rol van menselijk toezicht bij AI-systemen met een hoog risico, bijvoorbeeld voor medische diagnoses, ondermijnd.
  - D) Incorrect. Als patiënten medische dossiers mogen wijzigen op basis van AI-suggesties, kan dat leiden tot onnauwkeurigheden, is dit niet in lijn met gebruikelijke medische praktijken (waaronder professioneel toezicht vereist is) en levert dit juridische risico's op.

16 / 40

Een verzekeringsmaatschappij implementeert een nieuw kredietscoresysteem op basis van AI, met toegang tot zowel interne als openbare databases. De volgende risico's worden vastgesteld:

- **Een gebrek aan goede trainingsdata.** Als het model niet goed wordt getraind, is het lastig om mensen een nauwkeurige, eerlijke score te geven.
- **Integratie met andere toepassingen.** De AI-engine zal lastig te integreren zijn in de bestaande tamelijk complexe en op sommige punten verouderde toepassingsomgeving.
- **Niet-naleving van de AVG.** De Algemene Verordening Gegevensbescherming (AVG) bevat specifieke vereisten voor autonome verwerking van persoonsgegevens door geautomatiseerde systemen.
- **Transparantie en kwaliteit van het model.** Zowel medewerkers als klanten moeten de resultaten en beslissingen van het AI-model kunnen begrijpen.

De verzekeringsmaatschappij moet de AI Act naleven.

Welk risico is **niet** belangrijk voor naleving van de AI Act?

- A) Een gebrek aan goede trainingsdata
  - B) Integratie met andere toepassingen
  - C) Niet-naleving van de AVG
  - D) Transparantie en kwaliteit van het model
- 
- A) Incorrect. Om discriminatie of oneerlijke beslissingen te voorkomen moeten AI-systemen gebruikmaken van hoogwaardige data zonder bias. Slechte trainingsdata kunnen leiden tot bias of een onnauwkeurige kredietscore, wat een overtreding is van de vereisten in de AI Act.
  - B) Correct. Een gebrekkige integratie met andere toepassingen is absoluut zeker een risico, maar geen risico zoals gedefinieerd in de AI Act. (Literatuur: A, hoofdstuk 7.2, 7.3)
  - C) Incorrect. De AVG bevat specifieke beperkingen voor systemen waarin persoonsgegevens autonoom worden verwerkt, maar dat is niet de grootste uitdaging die moet worden aangepakt. De AI Act sluit aan op de AVG, vooral wat betreft de verwerking van persoonsgegevens, een rechtmatige basis voor AI-beslissingen en individuele rechten (bijvoorbeeld: het recht op uitleg en beroep).
  - D) Incorrect. De datakwaliteit en de nauwkeurigheid van het AI-model zijn de grootste uitdagingen die in dit soort toepassingsprojecten moeten worden aangepakt. Verder is het van wezenlijk belang dat de output van het AI-model begrijpelijk en verklaarbaar is. Onder de AI Act zijn verklaarbaarheid en transparantie vereist, vooral voor AI-systemen met een hoog risico, zoals hier voor het toekennen van kredietscores, omdat AI-beslissingen gevolgen hebben voor de financiële toegankelijkheid.

17 / 40

Een overheidsinstantie stelt voor een AI-systeem te gebruiken om brandhaarden voor criminaliteit in het centrum van een grotere stad beter te kunnen voorspellen. Het systeem zal worden gebruikt voor geautomatiseerde surveillance. Het wordt zodanig geprogrammeerd dat iedereen die verdacht gedrag vertoont automatisch wordt geïdentificeerd en aan de lokale politie wordt doorgegeven. Dit is een geweldige kans om misdaad te voorkomen, het veiligheidsgevoel te vergroten en te zorgen dat niemand zijn straf ontloopt.

Kleven er risico's aan de implementatie van dit AI-systeem?

- A) Ja, want een risico op bias is inherent aan AI-systemen die worden gebruikt voor geautomatiseerde beslissingen, waardoor personen mogelijk onterecht worden benadeeld.
  - B) Ja, want de AI Act voorziet zo veel privacyrisico's met bewakingssystemen dat de inzet ervan in openbare ruimtes ronduit verboden is onder deze wet.
  - C) Nee, want voor de vervolging en preventie van misdaad vormen AI-systemen geen specifieke risico's, aangezien ze worden ingezet om de openbare veiligheid te verbeteren.
  - D) Nee, want AI-systemen in het publieke domein vergroten de efficiëntie en vormen geen risico, aangezien alle beslissingen objectief en vrij van menselijke fouten zijn.
- 
- A) Correct. Deze belangrijke reden tot zorg wordt specifiek benoemd in de AI Act. Op gevoelige gebieden als misdaadbestrijding kunnen mogelijke bias in data of ondeugdelijke algoritmes in AI-systemen discriminerende resultaten opleveren. Voor AI-systemen in toepassingen met een hoog risico zijn onder de AI Act openheid, risico-evaluaties en bias-beperkende technieken vereist. (Literatuur: A, hoofdstuk 8.1, 8.2, 8.3)
  - B) Incorrect. De AI Act streeft ernaar veilig, open en eerlijk gebruik van AI te controleren en garanderen, in plaats van de inzet ervan in het publieke domein te ontmoedigen. Hoewel de AI Act beschermende maatregelen afdwingt om gevaren te ondervangen, stimuleert de wet creativiteit.
  - C) Incorrect. Verbetering van de openbare veiligheid is een nobel doel, maar neemt de verplichting niet weg om privacyrisico's en de benadeling van personen die zich in het openbaar niet misdragen zoveel mogelijk te beperken.
  - D) Incorrect. De output van AI is afhankelijk van de toegepaste trainingsdata en algoritmes, dus elke vorm van bias in de trainingsdata wordt overgeheveld naar de beslissingen die het systeem neemt. AI-beslissingen zijn niet noodzakelijkerwijs objectiever dan het menselijk beoordelingsvermogen. Bovendien geeft de AI Act personen het recht op menselijk toezicht op alle beslissingen die over hen worden genomen.

**18 / 40**

Een organisatie ontwikkelt een AI-systeem voor personeelswerving. Tijdens interne tests is het team op een risico gestuit: soms begunstigt het systeem onbedoeld sollicitanten met een bepaalde achtergrond, wat mogelijk tot discriminerende resultaten leidt. Het team twijfelt nu hoe het moet reageren op deze reden tot zorg.

De organisatie moet de AI Act naleven. Om het risico te beperken, hanteert de organisatie de norm ISO/IEC TR 24368.

Wat moet de organisatie volgens deze norm doen om dit risico te beperken?

- A)** Het algoritme aanpassen zodat het prioriteit geeft aan demografische quota's op basis van werkgelegenheidsstatistieken
  - B)** Overstappen op een dataloze aanpak door alle demografische data uit de trainingsset te verwijderen
  - C)** Cyberbeveiligingsmaatregelen toepassen om de data van sollicitanten te beschermen en de integriteit van het systeem te verbeteren
  - D)** Een proces voor de betrokkenheid van stakeholders implementeren om mogelijke bias op te sporen en weg te nemen
- 
- A)** Incorrect. Het klinkt misschien ethisch om te proberen tot een gelijkwaardige representatie te komen, maar het toepassen van harde quota's zonder context kan een nieuwe vorm van bias in de hand werken. ISO/IEC TR 24368 legt meer nadruk op billijkheid en betrokkenheid van stakeholders dan op willekeurige demografische targets.
  - B)** Incorrect. Door demografische data simpelweg te verwijderen wordt discriminatie niet voorkomen; dit kan zelfs bestaande vormen van bias minder zichtbaar maken. ISO/IEC TR 24368 stimuleert transparante methoden en het beperken van bias, niet het blind verwijderen van data.
  - C)** Incorrect. Cyberbeveiliging is weliswaar belangrijk, maar ethische kwesties als bias of billijkheid vallen hier niet onder. Voor de aanpak van ethische redenen tot zorg zijn methoden nodig die zijn afgestemd op ISO/IEC TR 24368.
  - D)** Correct. ISO/IEC TR 24368 bevordert de betrokkenheid van stakeholders om ethische risico's als bias aan het licht te brengen en inclusieve, eerlijke AI-systemen te ontwikkelen. Hiermee worden de doelstellingen wat betreft billijkheid en non-discriminatie uit de AI Act ondersteund. (Literatuur: B, hoofdstuk 4)

19 / 40

Een organisatie ontwikkelt een AI-systeem om leningen goed te keuren. Tijdens interne tests stuit het nalevingsteam op deze risico's: de manier waarop het model beslissingen neemt is onvoldoende transparant en de documentatie voor het evalueren van risico's is beperkt.

De organisatie moet de AI Act naleven en hanteert daarvoor de norm ISO/IEC 42001 en het NIST AI Risk Management Framework (RMF).

Wat moet het bedrijf volgens deze norm en dit kader doen om deze risico's aan te pakken?

- A) Een beoordeling van de veiligheidsnaleving uitvoeren op basis van aanbevolen cyberbeveiligingsrichtlijnen
  - B) Het AI-systeem onmiddellijk uit bedrijf nemen en overstappen op een handmatig goedkeuringsproces voor leningen
  - C) Een meetplan met transparantiewaarden definiëren en voor toezichtdoeleinden de logica achter beslissingen vastleggen
  - D) Het systeem opnieuw bouwen met synthetische data om zo veel mogelijk bronnen van bias weg te nemen
- 
- A) Incorrect. Het volgen van beveiligingsrichtlijnen is niet specifiek nodig om de transparantie of risicodocumentatie aan te pakken. Daarom is dit niet direct relevant voor de vastgestelde kwestie.
  - B) Incorrect. Het is niet nodig om het systeem uit bedrijf te nemen, want het bevindt zich nog maar in de interne testfase. Door over te stappen op handmatige goedkeuring kan het probleem inderdaad worden vermeden, maar dan zouden ook alle investeringskosten voor niets zijn geweest. De risico's kunnen worden beheerd met gestructureerde governance.
  - C) Correct. ISO/IEC 42001 legt de nadruk op transparante en verklaarbare besluitvorming, plus grondige documentatie binnen de volledige levenscyclus van AI. Het NIST AI RMF ondersteunt de definitie van meetwaarden aan de hand van de bijbehorende meetfunctie en stimuleert traceerbaarheid. Deze twee maatregelen samen leiden ertoe dat de kwesties in het scenario rechtstreeks worden aangepakt. (Literatuur: B, hoofdstuk 2.2, 2.5)
  - D) Incorrect. Enkel het gebruik van synthetische data biedt geen garantie dat bias wordt beperkt; bovendien worden zo fundamentele vereisten voor de transparantie en risicodocumentatie niet aangepakt.

20 / 40

Een gerenommeerde autofabrikant heeft een verregaand geautomatiseerd voertuig (niveau 4) ontwikkeld, uitgerust met een op AI gebaseerde technologie voor objectherkenning om de verkeersveiligheid te waarborgen. Tijdens het testen worden de volgende risico's vastgesteld:

- Bij weinig licht is het systeem onvoldoende in staat om verkeersdrempels te herkennen.
- Mogelijk zal het model lastig te verkopen zijn, omdat het de afmetingen van andere voertuigen niet weet.
- De ontwikkelaars kunnen niet precies uitleggen hoe het model beslissingen neemt.
- In de ontwikkelingsfase van het AI-systeem zijn niet alle stakeholders betrokken bij het ontwikkelproces.

Wat moet hier worden aangepakt als er **alleen** naar naleving van de AI Act wordt gekeken?

- A) Het risico op onvoldoende tests in realistische omstandigheden
  - B) Het risico op een gebrek aan transparantie in de besluitvorming door de AI
  - C) Het risico op een beperkte schaalbaarheid van het AI-systeem naar andere voertuigmodellen
  - D) Het risico op beperkte betrokkenheid van stakeholders tijdens de ontwikkeling van de AI
- 
- A) Incorrect. De AI Act is meer gericht op het beperken van vastgestelde risico's en het waarborgen van transparantie en grondrechten dan het aanpakken van onvoldoende realistische tests.
  - B) Correct. Artikel 11 van de AI Act benadrukt de transparantie van AI-systemen om mogelijke risico's op te sporen en verhelpen, waar het in dit scenario feitelijk om draait. (Literatuur: A, hoofdstuk 7.10)
  - C) Incorrect. Schaalbaarheid is commercieel gezien cruciaal, maar heeft geen directe invloed op de vereisten voor risicobeperking en transparantie uit de AI Act.
  - D) Incorrect. Hoewel betrokkenheid van stakeholders belangrijk is, is dat niet iets waar de AI Act zich op richt.

21 / 40

Een logistiek bedrijf werkt aan een AI-systeem om bezorgroutes te optimaliseren en brandstof te besparen. De organisatie overweegt twee mogelijkheden:

- Een closedsourcemodel van een leverancier die snellere installatie en geverifieerde nalevingscertificeringen garandeert.
- Een opensourcemodel dat meer aanpassingsmogelijkheden en transparantie biedt.

Het bedrijf moet de AI Act naleven, maar zoekt ook naar een goed evenwicht tussen innovatie en kosten.

Welk model past het **beste** bij dit bedrijf?

- A) Een closedsourcemodel, omdat zulke AI op zich veiliger is en meer wordt vertrouwd door instanties. Zo neemt de kans op niet-naleving af.
  - B) Een closedsourcemodel, omdat dit vooraf gecertificeerde naleving biedt voor AI. Hierdoor hoeft het bedrijf minder moeite te doen om te bewijzen dat het de AI Act naleeft.
  - C) Een opensourcemodel, omdat dit volledige transparantie voor de AI garandeert. Dat is gunstig voor de documentatie- en controleerbaarheidsvereisten.
  - D) Een opensourcemodel, omdat dit is uitgesloten van naleving van de AI Act. Deze uitsluiting komt doordat de broncode openbaar beschikbaar is.
- 
- A) Incorrect. Een closedsource-aanpak is niet automatisch meer conform de regels of veiliger.
  - B) Incorrect. Hoewel closedsourcemodelen nalevingscertificeringen kunnen omvatten, bieden ze mogelijk niet de benodigde flexibiliteit en openheid voor de behoeften van het bedrijf of veranderende wettelijke vereisten.
  - C) Correct. De volledige transparantie van opensourcemodelen sluit aan bij de criteria voor controleerbaarheid, traceerbaarheid en risicobeheer uit de AI Act. Dankzij deze voordelen kan het logistieke bedrijf makkelijker bewijzen dat het de wet naleeft. (Literatuur: A, hoofdstuk 6)
  - D) Incorrect. Opensourcemodelen zijn onder de AI Act niet vrijgesteld van nalevingsverantwoordelijkheden.

22 / 40

In de AI Act worden ethische principes voor AI-ontwikkeling gedefinieerd.

Wat is **geen** van deze principes?

- A) Verklaarbaarheid
  - B) Billijkheid
  - C) Verliespreventie
  - D) Respect voor AI-autonomie
- A) Incorrect. Dit is een principe in de AI Act. Het vereist dat AI-systemen transparant en begrijpelijk zijn, zodat gebruikers en stakeholders begrijpen hoe beslissingen worden genomen en welke redenering erachter zit.
- B) Incorrect. Dit is een principe in de AI Act. Het schrijft voor dat AI-systemen zodanig moeten worden ontwikkeld en ingezet dat ze zonder bias of discriminatie functioneren, zodat onpartijdige en rechtvaardige resultaten voor iedereen gewaarborgd zijn.
- C) Incorrect. Dit is een principe in de AI Act. Het benadrukt hoe belangrijk het is om AI-systemen te ontwerpen die risico's beperken en schade voorkomen, zodat gebruikers en alle personen die de gevolgen van AI-technologieën ondervinden, gegarandeerd veilig en beschermd zijn.
- D) Correct. Het juiste principe is respect voor menselijke autonomie. De AI Act richt zich in de eerste plaats op principes als billijkheid, verliespreventie en verklaarbaarheid, met het streven om te zorgen dat AI-systemen verantwoord, transparant en zonder bias worden ontwikkeld en gebruikt. (Literatuur: A, hoofdstuk 9.1)

23 / 40

Een start-up ontwikkelt een AI-systeem om gepersonaliseerd leren op scholen te ondersteunen door lesplannen af te stemmen op de behoeften van individuele leerlingen. Het systeem verzamelt gegevens over de prestaties en het leergedrag van leerlingen.

Wat moet de start-up onder de AI Act overwegen om een evenwicht tussen innovatie en regelgeving te vinden?

- A) Het systeem niet als 'hoog risico' labelen om bijkomende regeldruk te omzeilen en innovatie te stroomlijnen
  - B) Zorgen dat het AI-systeem een conformiteitsbeoordeling ondergaat en voldoet aan alle regelgeving voor systemen met een hoog risico
  - C) Robuuste functies voor databescherming implementeren maar gebruikersmeldingen achterwege laten om vertragingen tijdens de inzet te vermijden
  - D) Het systeem exclusief voor privéscholen in de handel brengen om de impact van nalegingsvereisten bij een hoog risico te beperken
- A) Incorrect. Het is onethisch om het systeem verkeerd te categoriseren om zo regelgeving te vermijden, en dit kan ernstige juridische gevolgen hebben.
- B) Correct. Uitvoeren van een conformiteitsbeoordeling en zorgen voor transparantie zijn cruciaal om alle regelgeving voor systemen met een hoog risico na te leven. (Literatuur: A, hoofdstuk 7.6; AI Act, artikel 6, bijlage III)
- C) Incorrect. Gebruikersmeldingen zijn essentieel voor de transparantie en naleving van regels voor databescherming.
- D) Incorrect. De naleving is niet noodzakelijkerwijs soepeler op privéscholen. Regels moeten worden gevolgd, ongeacht de markt.

24 / 40

De financiële instelling Fintegra implementeert een AI-systeem om frauduleuze transacties te detecteren. Voor analysedoeleinden moet het systeem toegang krijgen tot transactiegegevens en demografische gegevens van klanten. Fintegra moet voldoen aan het vereiste voor minimale gegevensverwerking uit de AI Act.

Hoe kan Fintegra het **beste** voldoen aan het vereiste voor minimale gegevensverwerking?

- A) De instelling moet alle transactiegegevens anonimiseren en alle data op basis waarvan de identiteit van natuurlijke personen kan worden herleid verwijderen, zelfs als die cruciaal is voor fraudedetectie.
  - B) De instelling moet alle persoonsgegevens, inclusief volledige namen en exacte adressen, verzamelen om de precieze analyse en de verbetering na verloop van tijd te waarborgen en moet de data zo lang als nodig veilig bewaren.
  - C) De instelling moet het verzamelen van gegevens beperken tot transactiegegevens die relevant zijn om fraude te detecteren, en vermijden dat persoonsgegevens zoals volledige namen of exacte adressen van klanten worden verwerkt.
  - D) De instelling mag verzamelde gegevens alleen delen met erkende leveranciers die voldoen aan de AI Act, zodat de interne verwerking van persoonsgegevens zoals volledige namen van klanten tot een minimum beperkt blijft.
- 
- A) Incorrect. Hoewel anonimisering belangrijk is, ondermijnt de verwijdering van cruciale gegevens voor fraudedetectie de effectiviteit van het AI-systeem. Bovendien is dit niet nodig onder het principe van minimale gegevensverwerking uit de AI Act.
  - B) Incorrect. Het verzamelen en opslaan van alle beschikbare gegevens is een overtreding van het principe van minimale gegevensverwerking, zelfs wanneer dit veilig gebeurt, en vergroot het risico op niet-naleving van de AI Act.
  - C) Correct. Het principe van minimale gegevensverwerking onder de AI Act vereist dat organisaties alleen gegevens verzamelen en verwerken die strikt noodzakelijk zijn voor het specifieke doel van het AI-systeem. Door de focus te leggen op transactiegegevens die relevant zijn voor fraudedetectie en alle onnodige persoonsgegevens verder te vermijden, voldoet de instelling aan dit vereiste. (Literatuur: A, hoofdstuk 4.1)
  - D) Incorrect. Dit is geen goede optie, aangezien het delen van gegevens met externe leveranciers mogelijk in strijd is met regels voor gegevensbescherming. Zelfs als Fintegra en de leverancier allebei voldoen aan de AI Act, is dit niet in lijn met het principe van minimale gegevensverwerking.

25 / 40

EduTech implementeert een adaptief leerplatform dat met behulp van AI leertrajecten voor leerlingen personaliseert. Het platform past het moeilijkheidsniveau van taken aan op basis van individuele prestaties.

Welk risico moet EduTech beperken om te waarborgen dat dit AI-systeem ethisch wordt gebruikt?

- A) Het risico op bias en discriminatie, want dit zou oneerlijke voor- of nadelen opleveren voor bepaalde leerlingen. Dit risico kan worden beperkt door datasets en algoritmes van het AI-systeem regelmatig te controleren en bij te werken.
  - B) Het risico op een bovenmatige afhankelijkheid van technologie, waardoor leerlingen mogelijk geen kritische denkvaardigheden ontwikkelen. Dit risico kan worden beperkt door het besluitvormingsproces van het AI-systeem vertrouwelijk te houden, zodat leerlingen worden gestimuleerd om meer zelf na te denken.
  - C) Het risico op inbreuken op de privacy, omdat gevoelige gegevens van leerlingen, waaronder hun prestaties, mogelijk verkeerd worden verwerkt of worden onthuld. Dit risico kan worden beperkt door meer aandacht te besteden aan betere technische prestaties van het AI-systeem.
  - D) Het risico op transparantieproblemen, omdat leerlingen en onderwijzers mogelijk niet begrijpen hoe beslissingen worden genomen. Dit risico kan worden beperkt door te zorgen dat het AI-systeem werkt zonder menselijk toezicht, om zo de billijkheid te waarborgen.
- 
- A) Correct. Bias en discriminatie vormen grote risico's in het onderwijs. Door datasets en algoritmes regelmatig te controleren en bij te werken kan het risico op bias en discriminatie worden beperkt. (Literatuur: A, hoofdstuk 7.6)
  - B) Incorrect. Transparantie is cruciaal voor ethisch AI-gebruik. Het stimuleren van kritisch denken is belangrijk in het onderwijs, maar dit heeft niets te maken met de transparantie van AI-systemen.
  - C) Incorrect. Met technische prestaties alleen kunnen ethische zorgen niet worden weggenomen en neemt het risico op inbreuken op de privacy evenmin af.
  - D) Incorrect. Hoewel beslissingen zonder menselijke tussenkomst de billijkheid kunnen bevorderen, vergroot het ontbreken van menselijk toezicht op AI het risico op bias en discriminatie. Menselijk toezicht is essentieel om ethisch gebruik te waarborgen.

26 / 40

Een ziekenhuisafdeling is gespecialiseerd in de diagnostiek en behandeling van ziektes. De afdeling ontwikkelt een diagnostisch AI-systeem dat moet helpen bij het opsporen van zeldzame aandoeningen. Het systeem analyseert patiëntgegevens, medische voorgeschiedenis en scans.

Het systeem wordt succesvol in gebruik genomen in de Verenigde Staten (VS). Enkele medisch specialisten in de Europese Unie (EU) willen het ook gaan gebruiken, maar begrijpen niet goed hoe het AI-systeem werkt. Bovendien beschikken ze niet over speciale kennis van en ervaring met AI-monitoring of het herkennen van storingen of foute diagnoses.

Wat is **geen** risico in verband met de overstap op dit AI-systeem?

- A) Het risico op ontbrekend effectief menselijk toezicht
  - B) Het risico op onjuiste diagnoses wegens bias in de automatisering
  - C) Het risico op wantrouwen wegens te weinig transparantie
  - D) Het risico op onbevoegde inzage in patiëntendossiers
- 
- A) Incorrect. Omdat het team dat het systeem wil gebruiken niet over de benodigde specialistische kennis beschikt, is het mogelijk niet in staat de AI te monitoren en storingen of onnauwkeurige output te herkennen.
  - B) Incorrect. Het ontbreken van specialistische kennis over een juiste interpretatie van de output kan leiden tot bias en foute diagnoses.
  - C) Incorrect. De medisch specialisten begrijpen niet hoe het AI-systeem werkt, wat kan leiden tot wantrouwen wegens gebrek aan transparantie.
  - D) Correct. Hoewel dataprivacy en onbevoegde inzage belangrijke aandachtspunten zijn, vormen ze in dit scenario geen specifiek risico voor de overstap op en het begrip van het diagnostische AI-systeem. (Literatuur: A, hoofdstuk 7.7)

27 / 40

Een bedrijf maakt een AI-systeem voor slimme-woningassistenten. Tijdens het testen ontdekt het team dat fouten bij de spraakherkenning, zoals verwarrende woorden die ongeveer hetzelfde klinken, onbedoelde acties tot gevolg hebben, bijvoorbeeld dat het verkeerde apparaat wordt ingeschakeld. Deze fouten kunnen leiden tot inbreuken op de privacy, zoals gespreksopnames zonder toestemming of een onjuiste gebruikersidentificatie, waardoor gevoelige informatie mogelijk met onbevoegde partijen wordt gedeeld.

De organisatie moet de AI Act naleven en gebruikt daarvoor het kader CEN/CLC/TR 18115.

Wat moet de AI-aanbieder volgens dit kader doen om het beschreven probleem aan te pakken?

- A) Een plan voor de betrokkenheid van stakeholders opstellen om verschillende standpunten over de werking van het AI-systeem te verkrijgen
  - B) Een ethische impactbeoordeling uitvoeren om inzicht te krijgen in de privacyrisico's van slimme-woningassistenten
  - C) De kwaliteit van trainingsdata verbeteren door gebruik te maken van systematische validatie en methoden voor foutcontrole
  - D) Door middel van versleuteling de dataveiligheid vergroten om spraakgegevens te beschermen en inbreuken te voorkomen
- 
- A) Incorrect. Het betrekken van stakeholders is gunstig om inzicht te krijgen in de bredere implicaties, maar daarmee wordt de acute technische behoefte aan een betere datakwaliteit niet aangepakt.
  - B) Incorrect. Hoewel ethische impactbeoordelingen belangrijk zijn, is daarmee de geringe datakwaliteit niet verholpen die de oorzaak vormt van de onjuiste interpretaties.
  - C) Correct. Een verbetering van de datakwaliteit door middel van systematische validatie en foutcontrole sluit aan bij het kader CEN/CLC/TR 18115, aangezien hiermee de behoefte aan nauwkeurige en volledige data wordt aangepakt om veilige en effectieve prestaties van het AI-systeem te waarborgen. (Literatuur: B, hoofdstuk 1)
  - D) Incorrect. Hoewel dataveiligheid belangrijk is, wordt door versleuteling het probleem met de datakwaliteit niet aangepakt, terwijl het daar in dit scenario wel om draait.

28 / 40

Een technologiebedrijf wordt betrapt op het gebruik van een AI-systeem dat in real time biometrische identificatie op afstand mogelijk maakt, wat nadrukkelijk verboden is onder de AI Act.

Wat is de juiste sanctie voor deze overtreding?

- A) Een formele waarschuwing zonder financiële sancties
  - B) Een administratieve geldboete tot € 7,5 miljoen of 1% van de totale wereldwijde jaarlijkse omzet voor het voorafgaande boekjaar
  - C) Een administratieve geldboete tot €15 miljoen of 3% van de totale wereldwijde jaarlijkse omzet voor het voorafgaande boekjaar
  - D) Een administratieve geldboete tot €35 miljoen of 7% van de totale wereldwijde jaarlijkse omzet voor het voorafgaande boekjaar
- 
- A) Incorrect. Een formele waarschuwing zonder financiële sanctie is niet voldoende voor een ernstige overtreding van de AI-regelgeving.
  - B) Incorrect. Deze boete is te laag voor een schending waarbij een bedrijf verboden handelingen verricht.
  - C) Incorrect. Hoewel deze boete fors is, sluit hij onvoldoende aan bij de ernst van de overtreding.
  - D) Correct. Deze sanctie sluit aan bij de maximaal mogelijke boete voor de ernstigste overtredingen onder de AI-regelgeving. (Literatuur: A, hoofdstuk 3.11; AI Act, artikel 52 en 99)

29 / 40

De AI Act legt extra nadruk op het belang van twee aspecten van AI-systemen: transparantie en traceerbaarheid.

Waarom zijn transparantie en traceerbaarheid belangrijk?

- A) Omdat ze cruciaal zijn om verantwoording te waarborgen en het vertrouwen in AI-systemen te stimuleren
  - B) Omdat dit verplichte vereisten zijn voor alle producten, inclusief AI-systemen
  - C) Omdat ze van wezenlijk belang zijn voor de betrouwbaarheid en automatisering van AI-systemen
  - D) Omdat de Europese, Chinese en Amerikaanse wetgeving deze aspecten gemeen hebben
- 
- A) Correct. Gedetailleerde informatie over gebruikte data helpt de beslissingen en acties van een AI-systeem te begrijpen, verklaren en doorgronden. Dankzij traceerbaarheid kunnen de besluitvormingsprocessen, datasets en systeemwerking van AI worden gecontroleerd en geauditeerd. Dit is cruciaal om bias, fouten en problemen met de verantwoording te herkennen. Transparantie en traceerbaarheid zijn belangrijk voor de verantwoording en het gebruikersvertrouwen in AI-systemen. (Literatuur: A, hoofdstuk 3.1)
  - B) Incorrect. Hoewel transparantie en traceerbaarheid belangrijk zijn voor AI-systemen, zijn ze niet verplicht voor alle producten.
  - C) Incorrect. Transparantie en traceerbaarheid zijn belangrijk voor de verantwoording en voor de mate waarin inwoners van de Europese Unie (EU) en gebruikers de AI-systemen en -technologieën vertrouwen (wat iets anders is dan de betrouwbaarheid ervan). Betrouwbaarheid en automatisering hangen meer samen met de prestaties en robuustheid van AI-systemen, en niet zozeer met deze twee principes.
  - D) Incorrect. Consistentie en homogeniteit tussen de drie belangrijkste wereldwijde AI-wetgevingen is geen aspect waar de Europese Unie (EU) rekening mee houdt. De AI Act is een Europese verordening. China en de Verenigde Staten hanteren andere aandachtspunten en wettelijke kaders.

30 / 40

Een retailorganisatie gebruikt een AI-systeem dat de manier waarop website-elementen worden weergegeven automatisch aanpast aan gebruikersvoorkeuren en het gebruikte apparaat. Het systeem geeft aanbevelingen voor producten en verbetert de gebruikerservaring op basis van de klikgeschiedenis en hoeveel tijd gebruikers op een pagina doorbrengen.

In welke categorie moet het gebruik van dit AI-systeem volgens de AI Act worden ingedeeld?

- A) Onaanvaardbaar risico
  - B) Hoog risico
  - C) Beperkt risico
  - D) Laag of geen risico
- 
- A) Incorrect. Het systeem voldoet niet aan de normen voor een onaanvaardbaar risico, waarbij AI-systemen de menselijke waardigheid, veiligheid of grondrechten in gevaar brengen. Het beïnvloeden van aankoopgedrag binnen een retailomgeving is op zich niet schadelijk.
  - B) Incorrect. AI-systemen die worden ingezet voor vakgebieden als de gezondheidszorg, het bankwezen of de werkgelegenheid, waar hoogstwaarschijnlijk wezenlijke rechten of veiligheidsrisico's in het spel zijn, worden als een hoog risico beschouwd. De manier waarop de hier beschreven technologie niet-gevoelige gegevens gebruikt om de vormgeving van een website te verbeteren, voldoet niet aan de criteria voor een hoog risico.
  - C) Incorrect. Hoewel het systeem keuzes van consumenten beïnvloedt, rechtvaardigen de bescheiden impact en het gebruik van niet-gevoelige gegevens eerder een indeling als laag of geen risico.
  - D) Correct. Het systeem gebruikt niet-gevoelige gegevens, draait in een omgeving waar het belang laag is, en is alleen van invloed op de aankoopervaring van gebruikers, dus wordt het geclassificeerd als laag of geen risico. (Literatuur: A, hoofdstuk 3.3, 3.4)

31 / 40

Een bedrijf ontwikkelt een AI-systeem voor geautomatiseerde besluitvorming in het eigen wervingsproces. Het AI-systeem screent cv's, kent een score toe aan sollicitanten en geeft aanbevelingen voor sollicitatiegesprekken.

Het bedrijf maakt zich zorgen dat er in het AI-systeem sprake is van bias die mogelijk het wervingsproces beïnvloedt.

Wat kan het bedrijf het **beste** doen om bias in het AI-systeem te beperken?

- A) Het AI-systeem toestaan dat het leert en zich aanpast zonder verdere menselijke tussenkomst
  - B) De bias in de trainingsdata negeren en focussen op de prestaties van het AI-systeem
  - C) Een divers ontwikkelteam samenstellen om het AI-systeem te maken en monitoren
  - D) Het AI-systeem trainen met één enkele databron om te zorgen voor consistentie
- 
- A) Incorrect. Als het bedrijf toestaat dat het AI-systeem leert zonder menselijke tussenkomst, kan dat leiden tot onbedoelde bias en een gebrek aan verantwoording. (AI Act, richtlijn 65)
  - B) Incorrect. Als bias wordt genegeerd, kan dat leiden tot discriminerende resultaten. Bovendien is dit in strijd met ethische richtlijnen. Het systeem moet leren om bias te herkennen en op basis daarvan aanpassingen te doen. (AI Act, richtlijn 72)
  - C) Correct. Een divers team kan helpen bias op te sporen en te beperken, zodat wordt gewaarborgd dat het AI-systeem billijk en inclusief is. (Literatuur: A, hoofdstuk 4.5; AI Act, richtlijn 81)
  - D) Incorrect. Het gebruik van één enkele databron kan het generaliserend vermogen van het AI-systeem beperken en juist bias introduceren. (AI Act, richtlijn 73)

**32 / 40**

Feline Finesse is een webshop met accessoires en kussens voor kattenbezitters, die onder andere gepersonaliseerde knuffels op basis van ingestuurde foto's kunnen bestellen. De webshop gebruikt een AI-systeem dat het volgende kan:

- Het past prijzen dynamisch aan op basis van de activiteit van consumenten.
- Het rangschikt zoekresultaten op basis van klantvoorkeuren.
- Het geeft persoonlijke aanbevelingen voor andere producten waarvan het denkt dat klanten die leuk vinden.

De webshop wijst klanten momenteel op alles wat het AI-systeem doet, en is zeer transparant over hoe het algoritme werkt. De CEO zet echter vraagtekens bij deze praktijk en wil weten welke mate van transparantie vereist is en hoe die de omzet beïnvloedt.

Wat moet de CEO weten over transparantie binnen de context van de AI Act?

- A) Met transparantie kan de webshop aan consumenten bewijzen dat het systeem objectief is. Consumenten hebben onder de AI Act het recht om te begrijpen hoe hun data wordt gebruikt, en dit begrip bevordert hun vertrouwen.
- B) Transparantie kan de grenzen of beperkingen van het AI-systeem zichtbaar maken. Consumenten verliezen mogelijk hun vertrouwen in het bedrijf zodat zij dit beseffen en dit is schadelijk voor de reputatie van het bedrijf.
- C) Transparantie is niet verplicht voor e-commerce. Personalisering vergroot het gebruiksgemak voor consumenten en verder hoeven zij niet te weten of begrijpen hoe het AI-systeem werkt.
- D) Transparantie blijft beperkt tot het beschikbaar stellen van de broncode van het AI-systeem. Het consumentenvertrouwen in het systeem neemt mogelijk af als consumenten begrijpen hoe het algoritme precies werkt.
- A) Correct. Transparantie is een pijler van de AI Act en een bepalende factor voor het publieke vertrouwen. (Literatuur: A, hoofdstuk 3.6)
- B) Incorrect. Hoewel het onthullen van beperkingen ook bij transparantie hoort, is het niet de bedoeling dat het vertrouwen daardoor afneemt. Het moet juist leiden tot meer vertrouwen, omdat hiermee verantwoordelijkheid en realistische verwachtingen worden gegarandeerd.
- C) Incorrect. Hoewel gebruiksgemak prettig is voor de meeste consumenten, schrijft de AI Act transparantie ook wettelijk voor. Steeds vaker zijn consumenten zich bewust van en maken zij zich zorgen over ethische AI-methoden. Daarom bevordert transparantie hun vertrouwen in het systeem.
- D) Incorrect. Transparantie beperkt zich niet tot het openbaar maken van broncode. Het omvat ook een duidelijke beschrijving van de operationele beleidsregels, het datagebruik en de besluitvorming van AI. Vertrouwen opbouwen en verantwoordelijkheid garanderen kan alleen als er sprake is van transparantie.

**33 / 40**

Voor een wervingsproces wordt een AI-systeem met een hoog risico gebruikt dat sollicitanten automatisch filtert op basis van hun kwalificaties. De gebruiksverantwoordelijke heeft echter geen mechanisme geïmplementeerd voor tussenkomst of toezicht door een mens in geval van twijfelachtige beslissingen.

Is menselijk toezicht vereist voor dit systeem onder de AI Act?

- A)** Ja, want menselijk toezicht is nodig voor tussenkomst in besluitvormingsprocessen.
  - B)** Ja, want menselijk toezicht waarborgt dat verplichtingen op het gebied van billijkheid en transparantie worden nageleefd.
  - C)** Nee, want geautomatiseerde systemen zijn zodanig ontworpen dat ze functioneren zonder menselijke tussenkomst.
  - D)** Nee, want wervingsprocessen brengen geen kritieke veiligheidsrisico's voor natuurlijke personen met zich mee.
- 
- A)** Correct. De AI Act benadrukt het belang van menselijk toezicht voor AI-systemen met een hoog risico om te waarborgen dat er een mechanisme voor menselijke tussenkomst is, vooral bij scenario's waarin beslissingen twijfelachtig kunnen zijn of aanzienlijke gevolgen voor betrokken personen kunnen hebben. (Literatuur: A, hoofdstuk 10.2.3)
  - B)** Incorrect. Hoewel billijkheid en transparantie belangrijke aspecten van de AI Act zijn, is menselijk toezicht alleen vereist in geval van een beperkt of hoog risico.
  - C)** Incorrect. De AI Act benadrukt het belang van menselijk toezicht voor AI-systemen met een hoog risico om te waarborgen dat er een mechanisme voor menselijke tussenkomst is. Dit geldt vooral bij scenario's waarin beslissingen twijfelachtig kunnen zijn of aanzienlijke gevolgen voor betrokken personen kunnen hebben.
  - D)** Incorrect. Hoewel het wervingsproces op zich geen veiligheidsrisico's met zich meebrengt, kijkt de AI Act naar de ethische en sociale implicaties van AI-systemen. Menselijk toezicht is vereist om zorgen op het gebied van billijkheid en transparantie aan te pakken. Dit zijn twee cruciale aspecten bij wervingsprocessen.

34 / 40

Een bedrijf treft voorbereidingen om een AI-model voor algemene doeleinden op de markt te brengen. Het model kan worden aangepast voor taken als automatisering van de klantenservice, contentcreatie en data-analyse. Het bedrijf is gevestigd buiten de Europese Unie (EU), maar is van plan het model in diverse EU-lidstaten te distribueren.

Wat is onder de AI Act **niet** vereist voordat het AI-model voor algemene doeleinden in de EU wordt gedistribueerd?

- A) Een gemachtigde in de EU aanstellen om alle nalevingskwesties af te handelen
  - B) De EU-regelgeving op het gebied van auteursrechten naleven om het model te trainen met auteursrechtelijk beschermde data
  - C) Een grondige audit uitvoeren om te verifiëren of alle EU-wetten en -regels volledig worden nageleefd
  - D) Een gedetailleerde samenvatting publiceren van de gebruikte content om het AI-model voor algemene doeleinden te trainen
- 
- A) Incorrect. Elke aanbieder die buiten de EU is gevestigd, moet een gemachtigde in de EU aanstellen om nalevingskwesties af te handelen. (AI Act, artikel 54)
  - B) Incorrect. Hoewel AI-modellen voor algemene doeleinden geen volledige conformiteitsbeoordeling hoeven te ondergaan, moeten ze wel voldoen aan de auteursrechtwetgeving van de EU, zodat beschermde content die bij het trainen wordt gebruikt, aan de wettelijke vereisten voldoet. (AI Act, artikel 53(1)(c))
  - C) Correct. Een volledige conformiteitsbeoordeling is alleen vereist voor AI-systemen met een hoog risico en AI-modellen voor algemene doeleinden vallen niet in die categorie. Daarom is het bedrijf niet verplicht om een volledige conformiteitsbeoordeling uit te voeren. (Literatuur: A, hoofdstuk 3)
  - D) Incorrect. Volgens de AI Act moet er een samenvatting worden gepubliceerd van de gebruikte data om het model te trainen. Zo wordt transparantie gewaarborgd. (AI Act, artikel 53)

**35 / 40**

Een organisatie zet een AI-systeem in voor predictief onderhoud van industriële installaties. Wanneer het systeem enkele maanden in gebruik is, begint het een zeer hoog aantal valse meldingen te genereren, waardoor de workflows ontregeld raken. Uit een onderzoek blijkt het volgende:

- De organisatie heeft geen rekening gehouden met dynamische omgevingsveranderingen op de werkvloer.
- De organisatie beschikt niet over een formeel proces om risico's na de uitrol opnieuw te beoordelen.

De organisatie moet de AI Act naleven. Om deze kwesties op te lossen, hanteert de organisatie de norm ISO/IEC 23894.

Wat moet de organisatie volgens deze norm doen om de kwesties aan te pakken?

- A)** Een workshop Human Centered Design (HCD) organiseren om de bruikbaarheid van het systeem te verbeteren
  - B)** Een proces voor risicomanagement met constante evaluatie en monitoring ontwikkelen
  - C)** Een audit van de cyberbeveiliging uitvoeren om mogelijke kwetsbaarheden op te sporen en aan te pakken
  - D)** Het AI-systeem vervangen door een eenvoudiger model op basis van regels dat makkelijker te beheersen is
- 
- A)** Incorrect. Hoewel HCD de bruikbaarheid verbetert, wordt hiermee de onderliggende oorzaak niet aangepakt. Deze oorzaak is dat er geen dynamisch en adaptief risicomanagement is voor ingezette AI-systemen.
  - B)** Correct. ISO/IEC 23894 benadrukt het belang van een geïntegreerd proces voor dynamisch en constant risicomanagement in de volledige levenscyclus van AI, ook na de uitrol. Bij een herevaluatie had het systeem kunnen worden aangepast voordat het hoge aantal valse meldingen werd gegenereerd. (Literatuur: B, hoofdstuk 3.2, 3.4)
  - C)** Incorrect. Het is onwaarschijnlijk dat de valse meldingen worden veroorzaakt door de cyberbeveiliging. Bij deze oplossing wordt voorbijgegaan aan het risicomanagement gedurende de hele levenscyclus en de noodzaak om het AI-systeem constant opnieuw te evalueren.
  - D)** Incorrect. Als de organisatie het systeem vervangt, wordt daarmee voorbijgegaan aan de nadruk die ISO/IEC 23894 legt op de iteratieve behandeling en herbeoordeling van risico's, wat de voorkeur heeft boven het simpelweg afschaffen van de technologie.

36 / 40

Een wagenparkbeheerder gebruikt een AI-systeem om het gedrag van bestuurders bij te houden en voorspellingen over benodigd onderhoud te doen. Het systeem verzamelt en verwerkt grote datavolumes, zoals gps-locaties, rijpatronen en indicatoren voor voertuigprestaties. Tijdens een recente audit bleek dat het bedrijf beschikte over onvoldoende procedures voor databescherming.

Onder de AI Act zijn databeheer en privacybescherming essentieel voor dit bedrijf.

Waarom is dat zo?

- A) Omdat het bedrijf daarmee prioriteit kan geven aan zakelijke doelstellingen en operationele efficiëntie
  - B) Omdat daarmee het gebruikersvertrouwen toeneemt, persoonsgegevens worden beschermd en onbevoegde inzage wordt voorkomen
  - C) Omdat ze beide verplicht zijn en het bedrijf juridische problemen en mogelijke sancties kan voorkomen door de AI Act na te leven
  - D) Omdat daarmee procedures voor dataverzameling worden gestroomlijnd, aangezien gebruikers niet langer toestemming hoeven te geven
- 
- A) Incorrect. Implementatie van databeheer en privacybescherming is niet bedoeld om het bedrijf te helpen meer prioriteit te geven aan efficiëntie dan aan naleving.
  - B) Correct. De AI Act benadrukt het belang van de bescherming van individuele privacy en ethisch databeheer, wat een nauwkeurige en eerlijke werking van AI-systemen ondersteunt. (Literatuur: A, hoofdstuk 4.3, 4.4, 4.6)
  - C) Incorrect. Hoewel naleving belangrijk is, is de AI Act primair gericht op de bescherming van individuele rechten en de handhaving van ethische normen.
  - D) Incorrect. Onder de AI Act en andere relevante regelgeving op het gebied van databescherming zijn gebruikerstoestemming en databescherming vereist. Het is onwettig en onethisch om deze vereisten te negeren.

37 / 40

Welk gebruik van een AI-systeem wordt onder de AI Act geclassificeerd als een beperkt risico?

- A) Een chatbot die is ontworpen om klanten te helpen met algemene vragen en die is geprogrammeerd om hen te laten weten dat het een AI-systeem is.
  - B) Een systeem voor gezichtsherkenning dat wordt gebruikt om klanten in openbare ruimtes in real time te identificeren, bijvoorbeeld in een overdekt winkelcentrum.
  - C) Een tool voor medische diagnoses om artsen te helpen om op basis van patiëntgegevens behandelingen aan te bevelen.
  - D) Een AI-systeem dat een autonoom voertuig bedient dat zonder menselijk toezicht op de openbare weg rijdt.
- 
- A) Correct. De AI Act categoriseert AI-systemen die met gebruikers interageren maar niet veel invloed op rechten, veiligheid of wettelijke verplichtingen hebben als een beperkt risico. Deze systemen moeten voldoen aan transparantieplichtingen, bijvoorbeeld door gebruikers te laten weten dat ze interageren met een AI-systeem. (Literatuur: A, hoofdstuk 3.3)
  - B) Incorrect. Afhankelijk van de genomen beslissingen na identificatie wordt dit systeem gecategoriseerd als een hoog risico of is het systeem mogelijk zelfs verboden vanwege de privacy- en surveillance-implicaties.
  - C) Incorrect. Deze tool valt onder een toepassing met een hoog risico, omdat het AI-systeem werkt met gezondheids- en veiligheidsgegevens.
  - D) Incorrect. Autonome voertuigen worden als een hoog risico beschouwd vanwege zorgen over de veiligheid en de impact van mogelijke ongelukken.

38 / 40

Een bedrijf ontwikkelt een AI-systeem voor het onderwijs. Het AI-systeem bepaalt of een leerling toegang krijgt tot materialen, wordt toegelaten tot een school of wordt ingedeeld in een klas. Het AI-systeem zal beschikbaar zijn via clouddiensten.

In welke categorie moet het gebruik van dit AI-systeem volgens de AI Act worden ingedeeld?

- A) Onaanvaardbaar risico
  - B) Hoog risico
  - C) Beperkt risico
  - D) Laag of geen risico
- 
- A) Incorrect. AI-systemen die aanzienlijke gevolgen voor natuurlijke personen kunnen hebben, bijvoorbeeld voor hun toegang tot onderwijs, worden onder de AI Act als een hoog risico geclassificeerd, niet als onaanvaardbaar risico, aangezien ze met strikte vereisten worden gereguleerd in plaats van rondit worden verboden.
  - B) Correct. AI-systemen die zijn ontworpen om toegang tot onderwijs te beoordelen, moeten als een hoog risico worden geclassificeerd, omdat ze aanzienlijke gevolgen voor natuurlijke personen kunnen hebben. De AI is rechtstreeks van invloed op de toegang tot onderwijsbronnen of de toelating van leerlingen, wat gevolgen heeft voor hun grondrechten. (Literatuur: A, hoofdstuk 3.3, 3.4; AI Act, artikel 6 (bijlage III))
  - C) Incorrect. Voorbeelden van AI met een beperkt risico zijn chatbots met AI-ondersteuning, aanbevelingssystemen of AI-assistenten die geen cruciale beslissingen nemen over rechten of mogelijkheden van mensen en hen niet kunnen uitsluiten van toegang tot onderwijs. Hier is het risico hoog.
  - D) Incorrect. Classificatie als een laag risico is geen nauwkeurige afspiegeling van de mogelijke risico's bij dit type AI-systeem, aangezien het vermogen om personen uit te sluiten van onderwijs grote gevolgen voor die personen kan hebben.

**39 / 40**

Een organisatie heeft een AI-systeem voor automatische personeelswerving ontwikkeld. Het team doet de volgende bevindingen tijdens het testen:

- Het systeem geeft sollicitanten met een bepaalde etnische achtergrond consistent een lagere score, omdat er sprake is van demografische bias in de trainingsdata.
- Er bestaat momenteel geen intern beoordelingsproces of feedbackmechanisme van relevante partijen dat op het risico op deze specifieke vorm van bias had kunnen wijzen.

De organisatie moet de AI Act naleven. Om deze kwesties op te lossen, hanteert de organisatie de norm ISO/IEC TR 24368.

Wat moet de organisatie volgens deze norm doen om de kwesties op te lossen?

- A) Een synthetische dataset maken om scheve demografische verhoudingen aan te pakken en de billijkheid te verbeteren
  - B) Transparantie implementeren om de verklaarbaarheid en verantwoording van het systeem te vergroten
  - C) De praktijken voor dataversleuteling aanscherpen en toegangsbeheer gebruiken om inbreuk te voorkomen
  - D) Een ethisch kader met input van stakeholders gebruiken om mensenrechtenkwesties te evalueren
- 
- A) Incorrect. Enkel het gebruik van synthetische data biedt geen garantie dat bias wordt beperkt. Bovendien worden zo de fundamentele vereisten voor ethische AI-ontwikkeling niet aangepakt. Daarnaast is dit geen onderdeel van de norm ISO/IEC TR 24368 waarnaar hier wordt verwezen.
  - B) Incorrect. Transparantie heeft niet direct betrekking op de billijkheid, ethische beoordelingsprocessen of verplichte betrokkenheid van stakeholders. De relevante oplossing om deze kwesties aan te pakken, is de implementatie van een ethisch beoordelingsproces.
  - C) Incorrect. Hoewel dataversleuteling en toegangsbeheer cruciaal zijn om informatiebeveiliging te waarborgen, is dat niet relevant voor de kwesties die hier spelen. Focus op de implementatie van een ethisch beoordelingsproces zou in dit geval relevanter zijn.
  - D) Correct. ISO/IEC TR 24368 benadrukt het belang van ethische kaders, praktijken met het oog op mensenrechten, de betrokkenheid van stakeholders en billijkheid bij de ontwikkeling van AI. Door een ethisch beoordelingsproces in het leven te roepen, kan discriminatie worden herkend en beperkt, in lijn met de basisprincipes van de norm. (Literatuur: B, hoofdstuk 4.2, 4.3)

40 / 40

Een start-up ontwikkelt een AI-model voor algemene doeleinden dat wordt getraind met openbaar beschikbare onlinecontent, waaronder nieuwsberichten, onderzoeksrapporten en posts op social media. Na de lancering ontvangt het bedrijf een juridische kennisgeving van een groep auteurs die beweren dat hun intellectuele eigendommen zonder toestemming zijn gebruikt om het model te trainen.

Wat moet er in dit geval gedaan worden om intellectuele eigendomsrechten te beschermen?

- A) Aanvoeren dat het AI-model voor algemene doeleinden kan worden beschouwd als open source en daarom vrijgesteld is van nalegingsverplichtingen op het gebied van auteursrechten
  - B) Beweren dat er sprake is van eerlijk gebruik onder de AI Act, aangezien de content openbaar beschikbaar was, en de dataset blijven gebruiken
  - C) De door AI gegenereerde output die sterke overeenkomsten met de betwiste werken vertoont verwijderen om claims in verband van schendingen van het auteursrecht te vermijden
  - D) Naleving waarborgen door details van de gebruikte trainingsdataset te documenteren en delen, inclusief de herkomst van de set
- 
- A) Incorrect. Opensourcemodelen zijn niet automatisch vrijgesteld van nalegingsverplichtingen op het gebied van auteursrechten als ze systeemrisico's vormen of te gelde worden gemaakt.
  - B) Incorrect. De AI Act bevat geen uitzondering voor eerlijk gebruik. Openbaar beschikbare content kan nog steeds auteursrechtelijk beschermd zijn.
  - C) Incorrect. De AI Act schrijft niet voor dat door AI gegenereerde output moet worden verwijderd puur omdat deze overeenkomsten vertoont met auteursrechtelijk beschermde werken.
  - D) Correct. Onder artikel 53 van de AI Act moeten aanbieders het trainingsproces documenteren en daarbij gedetailleerde informatie over de herkomst en kenmerken van de data vermelden. (Literatuur: A, hoofdstuk 3)

# Evaluatie

De juiste antwoorden op de vragen in dit voorbeeldexamen staan in onderstaande tabel.

| Vraag | Antwoord | Vraag | Antwoord |
|-------|----------|-------|----------|
| 1     | B        | 21    | C        |
| 2     | B        | 22    | D        |
| 3     | A        | 23    | B        |
| 4     | A        | 24    | C        |
| 5     | D        | 25    | A        |
| 6     | C        | 26    | D        |
| 7     | A        | 27    | C        |
| 8     | A        | 28    | D        |
| 9     | C        | 29    | A        |
| 10    | A        | 30    | D        |
| 11    | D        | 31    | C        |
| 12    | B        | 32    | A        |
| 13    | A        | 33    | A        |
| 14    | B        | 34    | C        |
| 15    | A        | 35    | B        |
| 16    | B        | 36    | B        |
| 17    | A        | 37    | A        |
| 18    | D        | 38    | B        |
| 19    | C        | 39    | D        |
| 20    | B        | 40    | D        |





EXIN



Certified for what's next

Contact EXIN

[www.exin.com](http://www.exin.com)