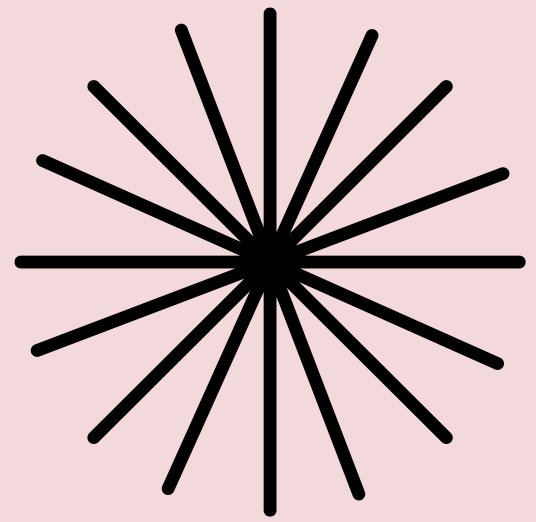
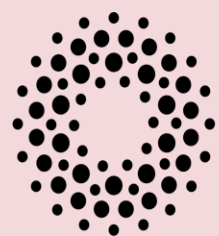




Risicobeoordelingsraamwerk voor AI-systemen -Concept-



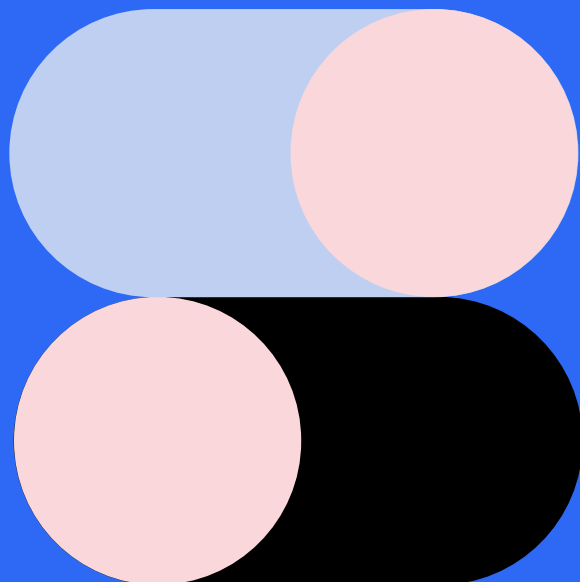
Npuls

4 september 2026

Conceptversie context

- Dit document bevat de conceptversie van het Risicobeoordelingsraamwerk voor AI-Systemen vanuit het project RAISE
- Het concept moet nog worden getest in de praktijk
- De laatste versie van dit raamwerk is steeds hier te vinden: <https://edu.nl/6huua>
- Neem voor vragen contact op met de projectgroep RAISE via aiendata@npuls.nl

Gebruik onder eigen verantwoordelijkheid



Op deze uitgave is de [Creative Commons](#) Naamsvermelding-GelijkDelen 4.0 Internationaal van toepassing.
Maak bij gebruik van dit werk vermelding van de volgende referentie: Risicobeoordelingsraamwerk voor AI-systemen - RAISE (2026). Utrecht: Npuls.

Dit moet je over het raamwerk weten

Dit risicobeoordelingsraamwerk helpt je grip te krijgen op de risico's van AI-systemen en impactvolle algoritmes in het onderwijs. AI-systemen vinden steeds sneller hun weg in het onderwijs. Dat biedt kansen, maar brengt ook risico's en verplichtingen uit wetgeving met zich mee.

Daarom heeft Npuls in co-creatie met onderwijsinstellingen uit het mbo, hbo en wo dit raamwerk ontwikkeld. Het raamwerk helpt je om:

- Snel te bepalen of, en zo ja hoe een systeem onder de AI-verordening valt.
- Inzicht te krijgen in het risiconiveau en de bijbehorende verplichtingen vanuit de wet.
- Risico's bij AI-systemen gestructureerd te beoordelen en te beheersen.

Het raamwerk is opgebouwd als een flexibele gids: je kunt het van begin tot eind doorlopen, maar ook gericht pagina's of hoofdstukken raadplegen op basis van je informatiebehoefte. Het raamwerk is geen gedetailleerd boekwerk dat antwoord geeft op alle mogelijke situaties. Daarentegen biedt het je een praktische en gestructureerde aanpak, die aansluit bij bestaande werkwijzen en helpt om compliant én verantwoord met AI om te gaan.

Afbakening

Het raamwerk is gericht op AI-systemen in het vervolgonderwijs (het mbo, hbo of wo). Het richt zich op onderwijs-specifieke toepassingen en AI-systemen binnen ondersteunende organisatieonderdelen, zoals de bedrijfsvoering. Het raamwerk gaat niet specifiek in op AI-toepassingen voor onderzoeksdoeleinden waarvoor aparte uitzonderingen in de AI-verordening gelden.

Definitie: In het raamwerk worden de termen 'AI-systeem' en 'AI-toepassing' door elkaar gebruikt. Daarmee wordt in beide gevallen een AI-systeem bedoeld volgens de definitie van de AI-verordening.

Inhoudsopgave



Waarom dit raamwerk?

Het raamwerk helpt je als onderwijsinstelling om te kunnen voldoen aan de AI-verordening en om gestructureerd risico's bij AI-systemen te beoordelen.

[Direct naar dit hoofdstuk](#)



Wat houdt het raamwerk in?

Overzicht van de stappen uit het raamwerk en uitleg hoe dit past binnen een bredere aanpak van het verantwoord gebruik van AI.

[Direct naar dit hoofdstuk](#)



Hoe pas ik het raamwerk toe?

Een begrijpelijke toelichting hoe je het raamwerk stap voor stap in de praktijk kunt toepassen.

[Direct naar dit hoofdstuk](#)

**Waarom dit
raamwerk?**

Doelstellingen van het raamwerk

AI-toepassingen worden in toenemende mate binnen onderwijsinstellingen in het vervolgonderwijs (mbo, hbo en wo) gebruikt. Dit biedt kansen voor het onderwijs, maar leidt ook tot nieuwe vragen en uitdagingen: hoe houd je als organisatie inzicht en overzicht op de verschillende systemen die je gebruikt en hoe houd je grip op de risico's? Dit is geen vrijblijvend vraagstuk; onderwijsinstellingen worden geconfronteerd met nieuwe wettelijke verplichtingen uit de AI-verordening. Daarom is er behoefte aan een praktische aanpak die in beide doelen ondersteunt.

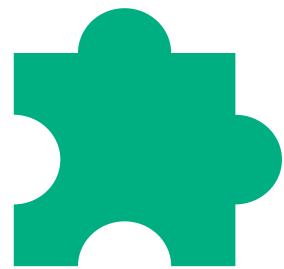
Daarom is Npuls samen met onderwijsinstellingen gestart met de ontwikkeling van dit **risicobeoordelingsraamwerk**. Hiermee worden twee samenhangende doelstellingen nagestreefd:

1. Inzicht krijgen in het risiconiveau van AI-systemen en de bijbehorende verplichtingen uit de AI-verordening
2. Het gestructureerd kunnen inschatten en beoordelen van risico's bij gebruik van AI-systemen

Op de volgende pagina worden de eisen en randvoorwaarden getoond op basis waarvan het raamwerk is ontwikkeld. Deze helpen je de keuzes achter het raamwerk beter te begrijpen.

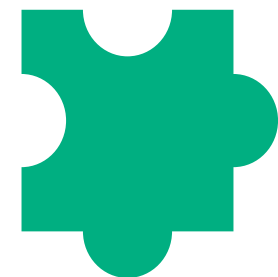
Eisen aan het raamwerk

Het raamwerk is ontwikkeld vanuit de volgende eisen en randvoorwaarden:



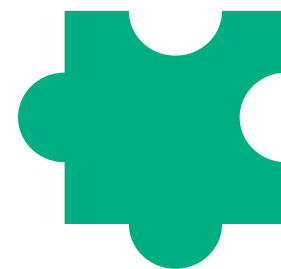
Praktisch bruikbaar

De praktische bruikbaarheid van het raamwerk staat voorop; daarom is gekozen voor compacte uitleg en beschrijvingen



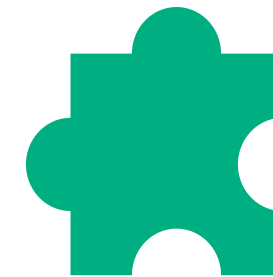
Niet uitputtend

De verplichtingen uit de AI-verordening vormen de leidraad; niet alle onderdelen om AI verantwoord in te zetten komen aan bod



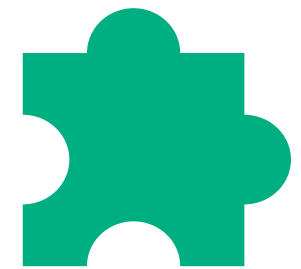
Inzet op hergebruik

Het doel is niet om het wiel opnieuw uit te vinden. Op verschillende plekken in het raamwerk wordt daarom (door)verwezen naar externe bronnen en hulpmiddelen



Standaard en flexibel

De stappen in het raamwerk zijn gestandaardiseerd, maar bieden ruimte voor maatwerk en flexibiliteit



Inpassing en integratie

Het raamwerk is zo ontworpen dat het makkelijk kan worden ingepast en geïntegreerd met huidige processen voor kwaliteit- en risicomanagement

Voor wie is dit raamwerk bedoeld?

- Het raamwerk gaat uit van het **Three Lines Model**. Het Three Lines Model is een manier om verantwoordelijkheden rond risico's en beheersing in organisaties te organiseren. We gaan ervan uit dat het raamwerk hoofdzakelijk wordt gebruikt door functionarissen in de tweede lijn die aan AI-compliance werken. Denk hierbij bijvoorbeeld aan: een AI officer (zowel compliance (AICO) als governance (ALGO)), (technical) information security officer ((T)ISO) of privacy officer (PO).
- Basiskennis van de inhoud van de AI-verordening wordt bij het gebruik van het raamwerk verondersteld – er wordt op meerdere plekken naar specifieke artikelen verwezen, maar de artikelen worden niet in detail toegelicht om het raamwerk compact en praktisch bruikbaar te houden.
- In de onderstaande tabel wordt het Three Lines Model verder toegelicht. Voor een goede risicobeoordeling kan input en betrokkenheid vanuit de eerste lijn natuurlijk niet ontbreken. De genoemde rollen in de tabel zijn indicatief; afhankelijk van de grootte en de governance-structuur binnen jouw instelling zullen de bevoegdheden van verschillende functies anders zijn.

Three Lines Model		
1e lijn – Uitvoering (business / operatie)	Verantwoordelijk voor het gebruik van AI-systemen en het beheersen van risico's in de praktijk	Rollen: bijvoorbeeld een docent, projectleider, proceseigenaar of eindgebruiker
2e lijn – Kaders en ondersteuning	Bepaalt de kaders voor risicobeheersing van AI-systemen en ondersteunt de 1e lijn bij de uitvoering	Rollen: bijvoorbeeld een AI officer (compliance en governance), privacy officer, (information) security officer, risicoadviseur of architect
3e lijn – Onafhankelijke controle	Toetst onafhankelijk of het geheel goed functioneert	Rollen: bijvoorbeeld een functionaris gegevensbescherming of (IT-)auditor

Gebruikers en andere stakeholders van het raamwerk

Gebruikers

Het raamwerk geeft dus vooral ondersteuning aan de samenwerking tussen de 1e en de 2e lijn binnen onderwijsinstellingen bij compliancy-inspanningen en risicobeheer van AI-systemen. Bij het hoofdstuk '[hoe pas ik het raamwerk toe](#)' wordt per stap toegelicht welke rollen vanuit het Three Lines Model logisch betrokken kunnen worden. Dit is geen exacte wetenschap; de daadwerkelijk te betrekken gebruikers en stakeholders zal per onderwijsinstelling verschillen. Zo heeft niet iedere onderwijsinstelling een aparte AI officer; soms wordt deze rol ingevuld door andere functies en rollen.

De meeste risicobeoordelingen kunnen effectief worden uitgevoerd in samenspel tussen 2e- en 1e-lijnsfuncties. De 2e lijn geeft de kaders aan en toetst de AI-systemen aan wet- en regelgeving en andere eisen. De 1e lijn participeert binnen dit proces door input te geven hoe, door wie en met welke gevolgen de AI-toepassing wordt gebruikt. Alleen door het samenspel tussen de 1e en 2e lijn kan een kwalitatief hoogstaande risicobeoordeling worden gedaan.

Andere stakeholders

Daarnaast kan, afhankelijk van het onderwerp, betrokkenheid en expertise van andere stakeholders nodig zijn. Met stakeholders bedoelen we rollen die niet altijd in de uitvoering betrokken zijn, maar op basis van het onderwerp en het vraagstuk betrokken kunnen worden. Denk hiervoor bijvoorbeeld aan studenten. Studenten krijgen vaak direct te maken met het toepassen van het AI-systeem of ondervinden effecten van het gebruik.

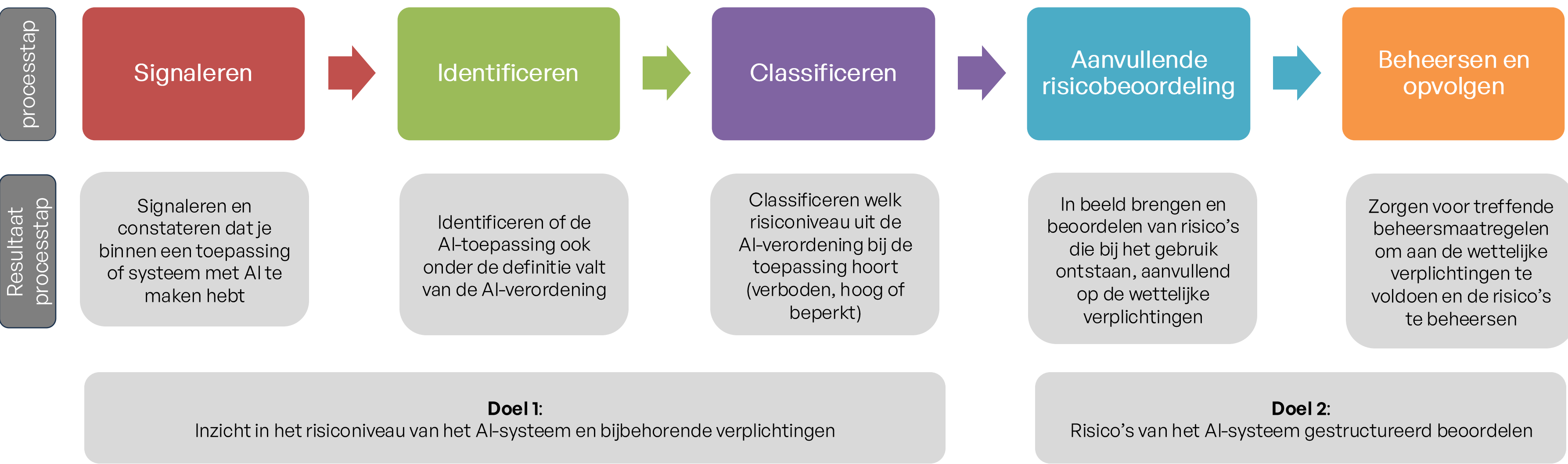
Om de relevante stakeholders in beeld te krijgen, is het raadzaam om een **stakeholderanalyse** te doen. Hieronder staat een onvolledige maar indicatieve lijst met voorbeelden van stakeholders die relevant kunnen zijn, afhankelijk van de type toepassing:

- Manager of bestuurder
- Datascientist
- Vertegenwoordiger van de leverancier van het AI-systeem
- Ethisch adviseur
- Inkoop- of contractadviseur
- Beleidsmedewerker
- Implementatieadviseur
- Jurist
- Onderwijsondersteunend personeel
- HR-medewerker
- Leden van de OR
- Student(vertegenwoordiging)

**Wat houdt het
raamwerk in?**

Overzicht van het raamwerk

Het raamwerk bestaat uit vijf processtappen. Het doorlopen van de eerste drie stappen voorziet in het eerste doel van het raamwerk: inzicht krijgen in het risiconiveau van het systeem en de bijbehorende verplichtingen uit de AI-verordening. Het doorlopen van de laatste twee processtappen zorgt ervoor dat je ook de risico's van AI-systemen gestructureerd kan beoordelen.



Wat is het raamwerk wél en wat is het niet?

Waar kan het raamwerk mee helpen? 	Wat is het raamwerk niet? 
Inzicht geven in het risiconiveau voor jouw AI-toepassing op basis van de definitie en eisen uit de AI-verordening	Het raamwerk alleen zorgt er niet voor dat je compliant bent met de AI-verordening. Iedere stap uit het raamwerk vraagt om concrete acties die je zelf als organisatie op de juiste manier moet invullen en uitvoeren om compliant te worden.
Een methode om binnen je onderwijsinstelling op een gestructureerde manier over risico's van AI-systemen na te denken en hier opvolging aan te geven	Aangeven welke risicobereidheid een onderwijsinstelling moet hebben bij het gebruik van AI-systemen
Het in de juiste volgorde doorlopen van stappen om een gestructureerde risicobeoordeling te kunnen doen	Een lineair proces dat altijd op dezelfde wijze wordt doorlopen, ongeacht de aard van toepassing die je gebruikt
Doorverwijzing naar achterliggende praktische hulpmiddelen die je helpen bij het verantwoord inzetten van AI	Een alles omvattend hulpmiddel om over de hele breedte van de organisatie AI op een verantwoorde manier in te zetten

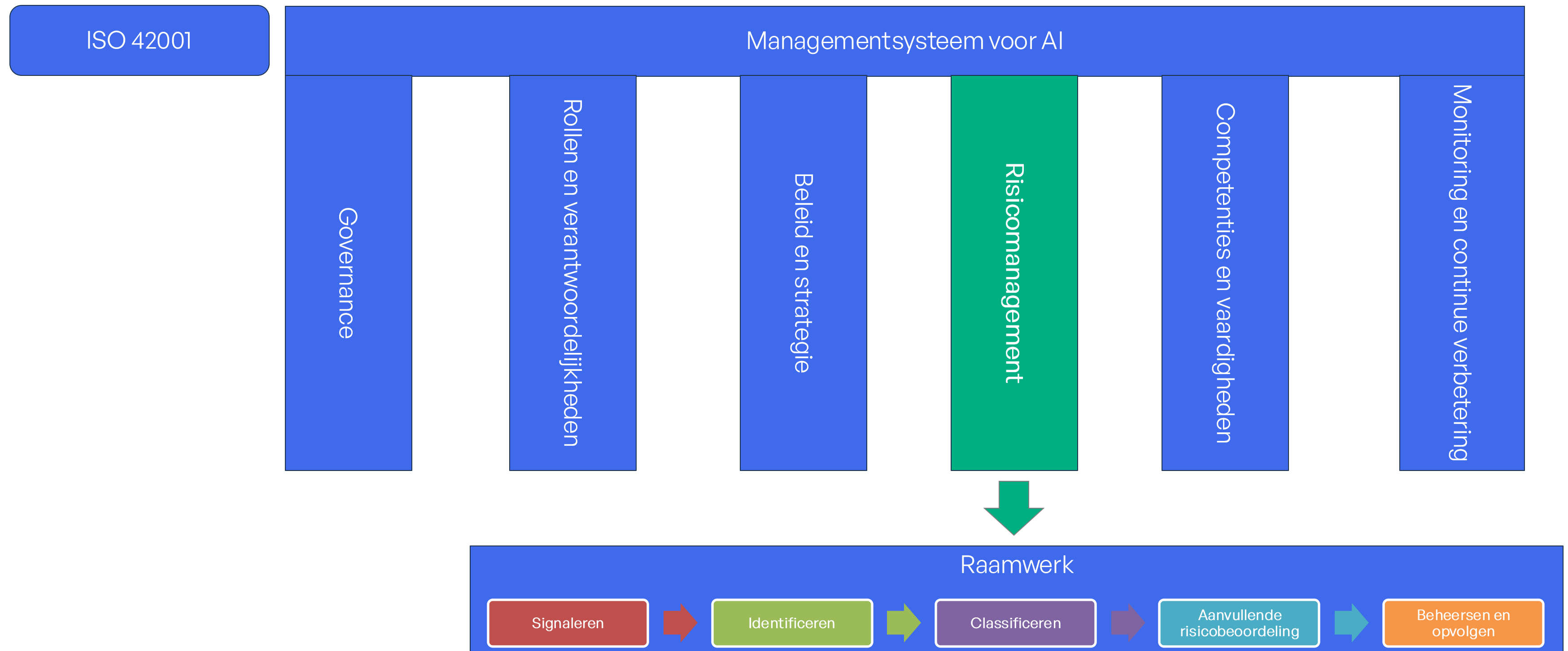
Hoe verhoudt het raamwerk zich tot ISO-normen?

- Binnen het vervolgonderwijs wordt er steeds meer gebruik gemaakt van ISO-normen om compliance-inspanningen te structureren, bijvoorbeeld op het gebied van informatiebeveiliging (ISO 27001). Ook voor het verantwoord inzetten van AI-systemen zijn er ISO-normen. De bekendste zijn:
 - **ISO 42001** biedt een structuur voor verantwoord AI-gebruik binnen organisaties. De norm specificeert de voorwaarden en vereisten en geeft ondersteuning bij het inrichten, implementeren, onderhouden en voortdurend verbeteren van een **AI-managementsysteem** binnen een organisatie.
 - **ISO 23894** beschrijft richtlijnen voor het identificeren, analyseren en mitigeren van AI-risico's. Deze norm sluit aan op de algemene ISO-norm voor risicomanagement (31000), maar kent een focus op specifieke risico's die bij AI-systemen aan de orde kunnen zijn, zoals bias en discriminatie en betrouwbaarheid van modellen.
- Dit raamwerk is geen 1-op-1 vertaling van deze ISO-normen, maar is wel in geest van de normen opgesteld. Het raamwerk is bovenal bedoeld om op een praktische wijze invulling te geven aan risicomanagement bij AI-systemen.

Merk op: Iedere onderwijsinstelling kan het raamwerk gebruiken. Er worden geen specifieke eisen aan je startsituatie gesteld. Houd er echter wel rekening mee dat niet alle elementen uit de ISO-normen voor verantwoord gebruik van AI in dit raamwerk aan de orde komen, terwijl er wel een onderlinge afhankelijkheid tussen deze elementen bestaat: gedegen risicomanagement bij AI-systemen kan alleen plaatsvinden als organisaties rollen en verantwoordelijkheden duidelijk hebben gescheiden en verdeeld. Dit geldt ook voor een organisatiebreed AI-beleid; dit is cruciaal om te kunnen bepalen of en zo ja welke AI-systemen je wilt inzetten, nog voordat je over risico's gaat nadenken. Het adresseren van deze elementen helpt daarom bij het effectief kunnen toepassen van het raamwerk.

Ter illustratie: relatie met ISO 42001

ISO 42001 biedt een holistisch managementsysteem voor verantwoord AI-gebruik. Dit raamwerk is géén manier om aan de hele ISO-norm te kunnen voldoen, maar richt zich op een specifiek onderdeel, geïllustreerd in het onderstaande figuur.



Ethisch aspect

Verantwoord gebruik van AI vraagt meer dan alleen compliance en risicomanagement. De uiteindelijke keuze om een AI-toepassing wél of niet te gebruiken hangt af van de risicobereidheid en is een normatieve beslissing van de onderwijsinstelling zelf. Naast de vraag wat mag volgens de wet en hoe risico's worden beheerst, vraagt het gebruik van AI nadrukkelijk om een expliciete afweging of een toepassing ook wenselijk is in het licht van publieke waarden, het eigen beleid en de missie van de onderwijsinstelling. Dit vraagt om een ethische afweging. Bij het lezen van het raamwerk zie je drie soorten vragen herhaaldelijk terugkomen: *Mogen* we de toepassing gebruik? *Willen* we de toepassing gebruiken? En zo ja, *hoe* zetten we de toepassing in?

Dit wordt in het onderstaande figuur nader toegelicht.

Vraag	Waar gaat het over?	Hoe het raamwerk hierin voorziet
<i>Mogen</i> we de toepassing gebruiken?	Valt de toepassing onder de AI-verordening en welke verplichtingen gelden?	Stap 2: Identificeren Stap 3: Classificeren
<i>Willen</i> we de toepassing gebruiken?	Past de toepassing bij de publieke waarden, beleid en missie van onze onderwijsinstelling? Dit is een strategische afweging (van het bestuur)	Het ethische aspect komt op verschillende plekken in het raamwerk terug en het meest expliciet bij Stap 1 (Signaleren) en tussen Stap 3 (Classificeren) en Stap 4 (Aanvullende risicobeoordeling). Er is geen vastomlijnde manier om een ethische afweging te maken, dit is aan de onderwijsinstelling zelf en hangt af van je eigen waarden, richtlijnen of procedures. Als je deze niet voor AI hebt, kan je inspiratie putten uit de Ethische Richtsnoeren van de Europese Commissie. Je kan ook gebruik maken van andere ethische hulpmiddelen, zoals de waardenwijzer van SURF & Kennisnet of de Impact Assessment Mensenrechten en Algoritmes (IAMA) van de Rijksoverheid.
<i>Hoe</i> zetten we de toepassing in?	Onder welke voorwaarden en met welke beheersmaatregelen zetten we de toepassing verantwoord in?	Stap 4: Aanvullende risicobeoordeling Stap 5: Beheersen en opvolgen

Relatie AI-verordening en (U)AVG

De AI-verordening en de (Uitvoeringswet-)AVG vullen elkaar aan. De AI-verordening richt zich op het reguleren van AI-systemen en op verplichtingen die afhankelijk zijn van het risiconiveau van die systemen. De (U)AVG blijft daarnaast onverkort van toepassing wanneer binnen een AI-systeem persoonsgegevens worden verwerkt. In de praktijk betekent dit dat een onderwijsinstelling beide juridische kaders in samenhang moet toepassen.

De AI-verordening beantwoordt primair de vraag welke specifieke verplichtingen gelden voor de ontwikkeling, inzet en het gebruik van AI-systemen. De (U)AVG stelt aanvullende eisen aan de verwerking van persoonsgegevens en waarborgt dat deze verwerking rechtmatig, noodzakelijk, proportioneel en transparant plaatsvindt, met passende bescherming van de rechten en vrijheden van betrokkenen. In dit raamwerk ligt de nadruk op het (kunnen) voldoen aan de verplichtingen uit de AI-verordening.

Een nadere beoordeling aan de hand van de (U)AVG is in ieder geval noodzakelijk wanneer sprake is van één of meer van de volgende situaties:

Er worden persoonsgegevens* van studenten, medewerkers of andere betrokkenen verwerkt.

De toepassing gebruikt leerdata, toetsresultaten, gedragsdata, stagegegevens of andere gegevens met een gevoelige context.

De toepassing ondersteunt of beïnvloedt beoordelingen, selecties, signaleringen of besluiten over personen.

De toepassing monitort gedrag, prestaties, aanwezigheid of interacties.

Er wordt gebruikgemaakt van grootschalige datasets, profilering, inferenties of koppelingen tussen databronnen.

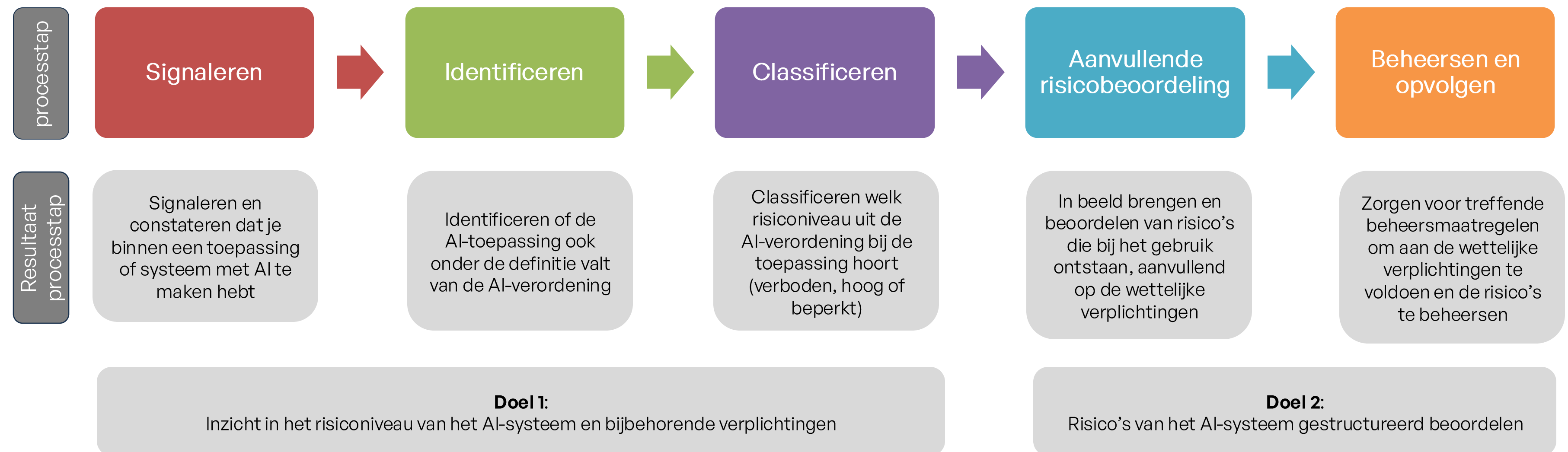
De uitkomst van het systeem kan significante gevolgen hebben voor de rechten, kansen of positie van betrokkenen.

* Alle informatie over een geïdentificeerde of identificeerbare natuurlijke persoon. Een natuurlijk persoon is een individueel mens. Het gaat zowel om gegevens die direct over iemand gaan als om gegevens die naar die persoon zijn te herleiden.

Hoe pas ik het
raamwerk toe?

Stappen van het raamwerk

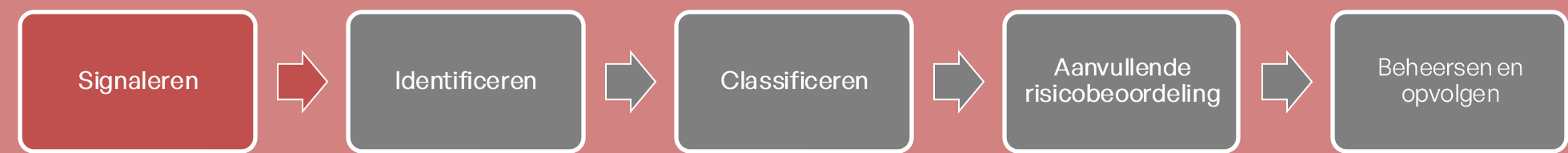
Op de volgende pagina's wordt per stap uitgelegd hoe je het raamwerk in de praktijk kunt toepassen. Daarbij staat een AI-systeem of AI-toepassing centraal. Met AI-toepassing bedoelen we AI-functionaliteit die onderdeel kan zijn van een groter systeem; de AI-verordening gebruikt hiervoor de term AI-systeem. Dit AI-systeem (of deze AI-toepassing) vormt in iedere stap het lijdend voorwerp dat centraal staat om te bepalen welke verplichtingen uit de AI-verordening van toepassing zijn en om op een gestructureerde manier een risicobeoordeling uit te voeren. Alle stappen in het raamwerk worden op dezelfde herkenbare manier beschreven. Op de volgende pagina wordt hiervan een voorbeeld weergegeven. Voor vastlegging van de uitkomst van de verschillende stappen kan je [een use case card](#) gebruiken.



Stap 0: Voorbeeld



Stap 1: Signaleren



WAAROM?

De eerste stap van het raamwerk, signaleren, heeft als doel om vast te stellen of er AI binnen een ICT-toepassing of systeem wordt gebruikt. Dit lijkt evident, maar is dat niet altijd; van veel toepassingen wordt snel geroepen dat er sprake is van AI (als paraplu-term), terwijl dit niet zo hoeft te zijn. Ook worden toepassingen gebruikt die misschien geen AI lijken, maar dit wel zijn. Het tijdig herkennen van AI-gebruik binnen de onderwijsinstelling zet het raamwerk in gang en maakt een passende risicobeoordeling mogelijk.



MET WIE?

- **De eerste lijn** sluit aan als initiatiefnemer en is verantwoordelijk voor het beschrijven van het doel, de context en beoogde gebruik van het AI-systeem. **De tweede lijn** sluit aan voor de compliance-begeleiding.
- **Mogelijke stakeholders zijn:** een ethisch adviseur, beleidsmedewerker, manager en/of bestuurder, vooral om te beoordelen en af te wegen of het voorgenomen AI-systeem aan het beleid en de waarden van de onderwijsinstelling voldoet en wenselijk is.



HOE?

De stap signaleren bestaat uit 3 onderdelen:

- Het **daadwerkelijk signaleren** van een potentiële AI-toepassing binnen je instelling; word er mogelijk AI in een ICT-toepassing of systeem gebruikt?
- **De context** van deze AI-toepassing in beeld brengen; hoe wordt AI binnen de toepassing gebruikt?
- Nagaan of de toepassing in lijn is met het beleid, de doelen en kaders die er binnen je instelling voor AI zijn: **willen we dit eigenlijk wel?**



RESULTAAT:

Na deze stap is het gebruik van een (mogelijke) AI-toepassing binnen de onderwijsinstelling gesignaleerd en is de context waarin de AI-toepassing wordt gebruikt duidelijk. Daarnaast is een ethische afweging van de wenselijkheid van de toepassing gemaakt. Als je beslist om niet met het inkopen, bouwen of gebruiken van de AI-toepassing door te gaan, dan hoef je de rest van het raamwerk niet te doorlopen. Als je wel doorgaat, ga je verder met Stap 2 (Identificeren) waarin je nagaat of de toepassing onder de definitie van de AI-verordening valt.

Nadere toelichting bij Signaleren

Signaleren van een AI-toepassing

Het signaleren van een AI-systeem binnen je instelling kan lastig zijn. Het is daarom behulpzaam om hier een duidelijke en concrete vorm aan te geven, bijvoorbeeld door het opzetten van een meldpunt. Hulpvragen die kunnen helpen bij het signaleren van een AI-systeem zijn:

- Welke taken worden (deels) geautomatiseerd door het gebruik van digitale tools?
- Wordt er gebruik gemaakt van tools die beslissingen nemen of voorspellingen doen? Of tools die (persoonlijke) aanbevelingen doen voor studenten of docenten?
- Zijn er afdelingen, medewerkers of studenten die zelf digitale tools ontwikkelen?
- Zijn er systemen die grote hoeveelheden data verwerken?
- Zijn er plannen om nieuwe systemen aan te schaffen die mogelijk AI-toepassingen bevatten?

De context van de AI-toepassing

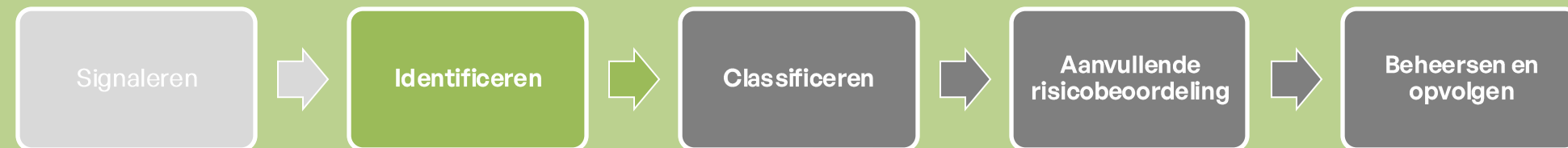
Voor je verder gaat met het raamwerk, is het belangrijk om de context van de gesignaleerde AI-toepassing duidelijk in beeld te brengen. Dit helpt bij het doorlopen van de volgende stappen in het raamwerk. Ontwikkel daarnaast een beeld van de stakeholders die je wil betrekken, zodat deze op de juiste momenten uitgenodigd kunnen worden. Hulpvragen voor het bepalen van de context zijn:

- Welk probleem wordt opgelost of welke kans wordt benut door het gebruik van deze toepassing?
- Wie maken er gebruik van deze toepassing? En wie krijgen te maken met de effecten / uitkomsten?
- Wat is het beoogde gebruik van deze toepassing?
- Is de toepassing al in gebruik, gaan we de toepassing aanschaffen of is de toepassing (intern) nog in ontwikkeling?

Willen we dit eigenlijk wel?

Als onderdeel van de stap signaleren is het belangrijk om binnen je instelling na te gaan of de toepassing in lijn is met de doelen en kaders die je voor AI hebt. Als dit immers niet zo is, dan stop je met het gebruiken van de toepassing en hoef je de verdere stappen in het raamwerk niet te doorlopen. Je stelt dus de vraag of de toepassing bij het beleid, de strategische doelen en de publieke waarden van je instelling past. Maak voor het maken van de ethische afweging gebruik van bestaande ethische hulpmiddelen, zoals de Ethische Richtsnoeren van de Europese Commissie, de waardenwijzer van SURF & Kennisnet of IAMA van de Rijksoverheid.

Stap 2: Identificeren



WAAROM?

Het doel van deze stap is om te identificeren in hoeverre een gesignaleerde AI-toepassing onder de definitie van de AI-verordening valt. Als dit het geval is, dan moet het risiconiveau van de AI-toepassing worden bepaald (zie [stap 3: classificeren](#)). Als de toepassing niet onder de AI-verordening valt, bijvoorbeeld omdat er sprake is van een regelgebaseerd algoritme, dan kan stap 3 uit het raamwerk worden overgeslagen. Kijk op [pagina 23](#) voor een schematische weergave van de twee mogelijke uitkomsten van deze stap.



MET WIE?

- **De eerste lijn** is verantwoordelijk om de (technische) werking van het systeem te doorgronden, al dan niet op basis van informatie van de leverancier of ontwikkelaar. **De tweede lijn** is vanuit de compliance-begeleiding betrokken om de definitie van AI-systeem uit de AI-verordening te interpreteren.
- **Mogelijke stakeholders zijn:** een data-scientist, de leverancier/ontwikkelaar om de benodigde achtergrondinformatie te geven of een jurist om te ondersteunen bij een juridische interpretatie van de wet.



HOE?

1. Gebruik een identificatietool om te bepalen of de toepassing onder de definitie van AI in de AI-verordening valt, bijvoorbeeld [het toetsingskader van Sivon](#), de [identificatietool van Algorithm Audit](#) of de [beslishulp van de Rijksoverheid](#) (toelichting: [pagina 22](#)).
2. Vul de tool bij voorkeur gezamenlijk tussen de 1e en 2e lijn in, zodat verschillende expertises worden meegenomen.
3. Bespreek de resultaten om tot een gedeeld oordeel te komen.



RESULTAAT:

Na deze stap is duidelijk of de toepassing onder de definitie van AI van de AI-verordening valt. Als dit het geval is, dan ga je door met de volgende stap, namelijk Stap 3 Classificeren. Zo niet, bijvoorbeeld omdat er sprake is van een regelgebaseerd algoritme, dan hoeft stap 3 niet te worden uitgevoerd. Je kan dan direct door met het uitvoeren van [stap 4: aanvullende risicobeoordeling](#).

Toelichting verschillende tools

Op dit moment is er nog geen specifieke tool ontwikkeld voor het vervolgonderwijs, maar de AI-verordening is dusdanig generiek dat andere tools ook goed bruikbaar kunnen zijn. Hieronder drie suggesties van tools die je kan gebruiken:

Het toetsingskader van Sivon

Dit toetsingskader biedt handvatten om AI-toepassingen in het onderwijs te beoordelen op basis van de AI-verordening, zowel identificatie als classificatie.

Let op: het is geen doorklikbare tool, maar een ondersteunend kader met criteria en toelichtingen. Hoewel het kader zich specifiek richt op het funderend onderwijs (basis- en voorgezet onderwijs) en met name bedoeld is voor leveranciers, bevat het waardevolle onderwijsspecifieke inzichten die ook voor het vervolgonderwijs nuttig zijn.

Tool van Algorithm Audit

Deze tool ondersteunt bij de identificatie en classificatie van AI-toepassingen op basis van de AI-verordening. Via interactieve vragenlijsten helpt de tool organisaties bij het identificeren van AI-systemen, het bepalen van je rol, risiconiveau en bijbehorende verplichtingen.

Let op: de tool is generiek ontwikkeld en niet specifiek gericht op het onderwijs, maar bevat wel een aantal onderwijsspecifieke vragen die zijn afgeleid van de AI-verordening.

Beslischulp van de Rijksoverheid

Deze beslischulp helpt organisaties bij de identificatie en classificatie van AI-toepassingen op basis van de AI-verordening en de daarbij behorende vereisten. Via een stapsgewijs proces worden criteria doorlopen, waarbij resultaten inclusief notities kunnen worden vastgelegd in een rapport.

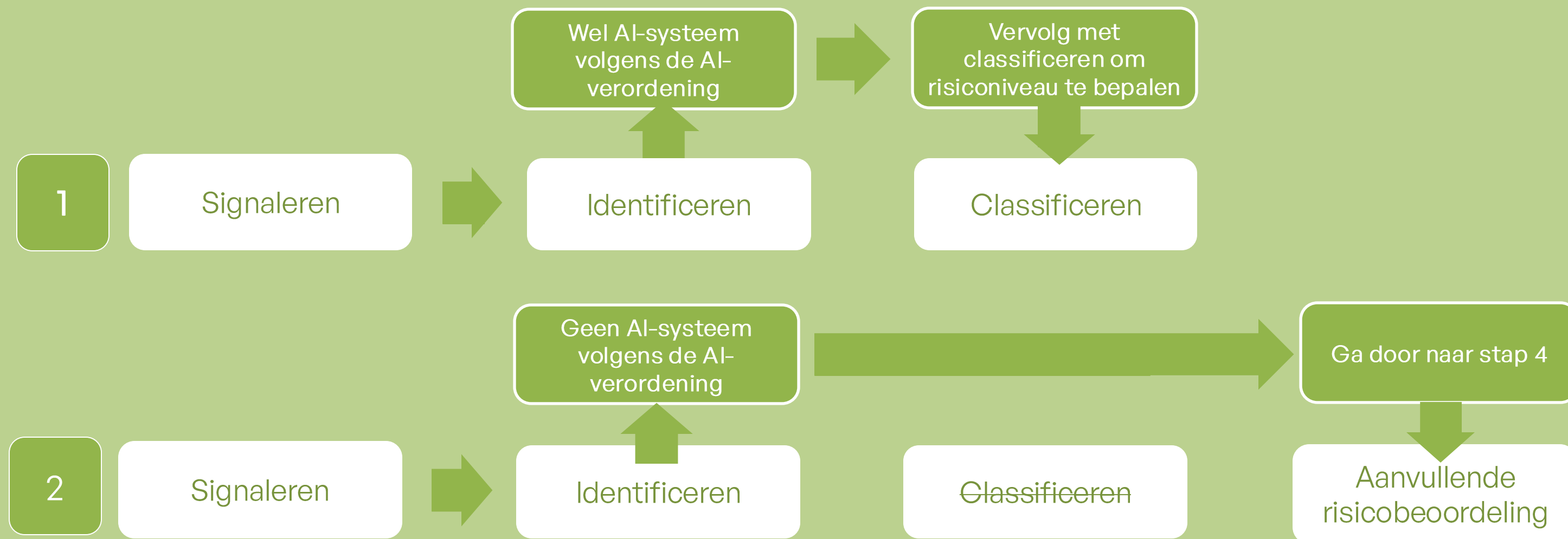
Let op: de tool is ontwikkeld voor overheidsorganisaties en niet specifiek voor het onderwijs, maar biedt wel een praktische manier om overzicht te creëren.

Identificatie kan leiden tot twee mogelijke uitkomsten

Stap 2 'identificeren' kan resulteren in twee mogelijke uitkomsten:

- (1) de toepassing waarvoor je het raamwerk doorloopt voldoet wél aan de definitie van de AI-verordening, óf
- (2) de toepassing waarvoor je het raamwerk doorloopt voldoet niet aan de definitie van de AI-verordening.

Zie onderstaande figuur wat de uitkomst voor de rest van het doorlopen van het raamwerk betekent.



Stap 3: Classificeren



WAAROM?

Het doel van deze stap is om vast te stellen welke rol je hebt (aanbieder of gebruiksverantwoordelijke) en in welk risiconiveau jouw AI-systeem valt op basis van de AI-verordening:

- Onaanvaardbaar risico = direct stoppen door de AI-toepassing niet in gebruik te nemen of uit gebruik te nemen.
- Hoog risico = voer een aanvullende risicobeoordeling uit; dit is wettelijk verplicht.
- Beperkt risico = de aanvullende risicobeoordeling is een eigen keuze, die o.a. afhangt van het eigen beleid.

Zie het figuur op [pagina 25](#) voor een schematische weergave.

MET WIE?

- **De eerste lijn** is verantwoordelijk om de context over de toepassing, het doel en de impact op processen van het AI-systeem in kaart te brengen. **De tweede lijn** is vanuit de compliance-begeleiding betrokken om de risiconiveaus uit de AI-verordening en de bijbehorende verplichtingen te interpreteren.
- **Mogelijke stakeholders zijn:** een jurist voor benodigde achtergrondinformatie, en/of een implementatieadviseur voor hernieuwde afweging over uitvoerbaarheid van de compliance-opgave.

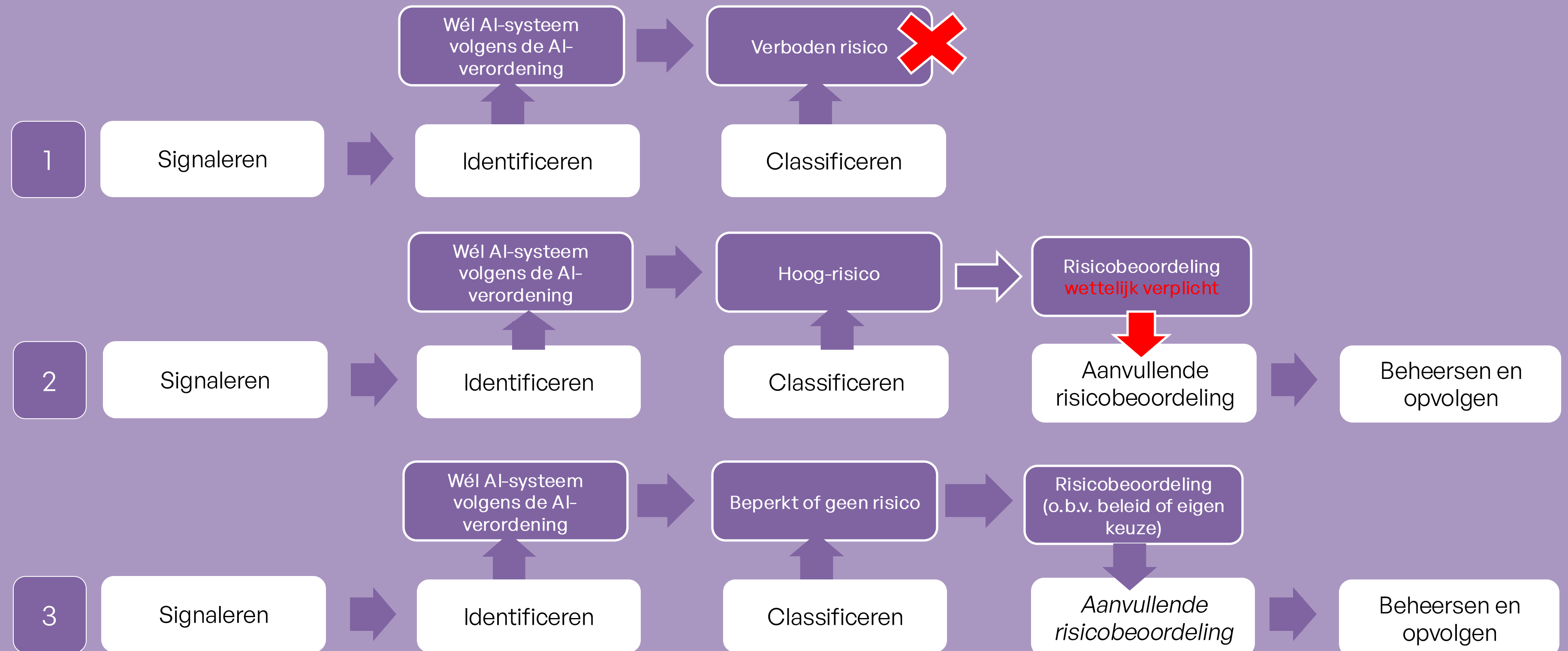
HOE?

1. Bepaal of je de rol van aanbieder of gebruiksverantwoordelijke hebt.
2. Gebruik een classificatietool om vast te stellen welk risiconiveau jouw toepassing kent, bijvoorbeeld het [toetsingskader van Sivon](#), de [identificatietool van Algorithm Audit](#) of de [beslischulp van de Rijksoverheid](#).
3. Vul de tool bij voorkeur gezamenlijk tussen de 1e en 2e lijn in, zodat verschillende expertises worden meegenomen.
4. Bespreek de resultaten om tot een gedeeld oordeel te komen.

RESULTAAT:

Na het uitvoeren van deze stap is het duidelijk onder welk risiconiveau het AI-systeem valt op basis van de AI-verordening (verboden, hoog of beperkt). De uitkomst van deze stap bepaalt het vervolg binnen het raamwerk. Als er sprake is van een verboden toepassing, moet het gebruik worden gestaakt. Als er sprake is van hoog-risico dan gelden er diverse wettelijke verplichtingen. Een overzicht hiervan staat op de volgende pagina's uitgelicht. Als je weet aan welke verplichtingen je moet voldoen, kan je een hernieuwde afweging maken: *willen* we met deze toepassing doorgaan? Kunnen we als instelling deze compliance-last dragen en zorgen voor verantwoord gebruik?

Classificatie bepaalt of een aanvullende risicobeoordeling verplicht of gewenst is



Verplichtingen per risiconiveau

Stap 3 ‘classificeren’ kan resulteren in drie mogelijke risiconiveaus (zie hieronder). Als het systeem niet onder een van de risiconiveaus valt, geldt alleen de generieke verplichting van AI-geletterdheid. Ieder risiconiveau kent aparte verplichtingen op basis van de AI-verordening, waarbij geldt: hoe hoger het risiconiveau, hoe zwaarder de verplichtingen. Op de volgende pagina staat een overzicht van de verschillende verplichtingen per risiconiveau en rol.

- **Verboden AI-systeem** (onaanvaardbaar risico) betekent niet doorgaan met het gebruiken van de AI-toepassing.
- **Hoog-risico AI-systeem** betekent verschillende verplichtingen, zowel voor gebruiksverantwoordelijken als aanbieders. Voor sommige hoog-risico toepassingen gelden er uitzonderingen. De uitzonderingsgronden zijn te vinden in artikel 6 (lid 2-6) van de AI-verordening. Hierin is bepaald dat systemen niet hoog-risico zijn als ze geen significante risico's voor gezondheid, veiligheid of grondrechten opleveren, bijvoorbeeld vanwege hun ondersteunende karakter of omdat ze alleen een voorbereidende functie hebben.
- **Beperkt risico AI-systeem:** transparantieverplichtingen.
- **Algemene verplichtingen:** als de toepassing niet onder een van de risiconiveaus valt, geldt de generieke verplichting van AI-geletterdheid.
→ Soms is het lastig te interpreteren in welk risiconiveau je toepassing valt. Gebruik daarom een classificatietool en/of kijk naar de zogeheten richtlijnen (*guidelines*) die de Europese Commissie publiceert met toelichting bij de verschillende risiconiveaus.

De wet stelt meer en zwaardere verplichtingen voor aanbieders dan voor gebruiksverantwoordelijken. Organisaties die zelf een AI-toepassing ontwikkelen of ingrijpend aanpassen, worden als aanbieder aangemerkt. Omdat de AI-verordening productregelgeving is, brengt deze rol belangrijke aanvullende verplichtingen met zich mee. Een bewuste afweging van deze rol is daarom essentieel.

Na deze stap beschik je over meer informatie over de wettelijke verplichtingen die je hebt bij jouw AI-toepassing. Dit vormt een hernieuwd moment om af te wegen of je dit als onderwijsinstelling wilt aangaan. Niet alleen ethisch (past het bij onze waarden en beleid?), maar: willen en kunnen we het risiconiveau aan? Hebben we voldoende compliance-capaciteit?

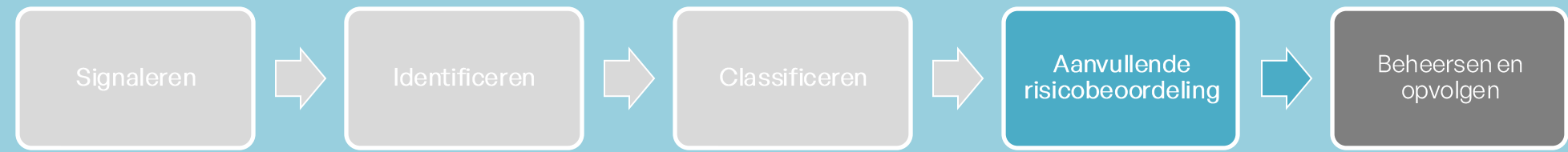


Verplichtingen uit de AI-verordening (2024) bij risiconiveaus

Risiconiveau	Verplichtingen aanbieder	Verplichtingen gebruiksverantwoordelijke
Onaanvaardbaar (=verboden) Zie: Art. 5	Verbod op in de handel brengen en gebruik	Verbod op gebruik
Hoog-risico Zie: Art. 6 en Bijlage I en III	Aanbieders moeten zorgen dat hun AI-systemen met hoog risico in overeenstemming zijn met de volgende eisen (Art. 16 lid a): - Een risicomanagementsysteem wordt opgezet, uitgevoerd, gedocumenteerd en onderhouden (Art. 9) - Trainings-, validatie- en testdatasets voldoen aan relevante kwaliteitseisen en governancevereisten (Art. 10) - Technische documentatie wordt opgesteld conform Bijlage IV (Art. 11) - Loggingfunctionaliteiten aanwezig zijn voor automatische registratie van gebeurtenissen gedurende de levenscyclus van het systeem (Art. 12) - Het systeem voldoende transparant is en gebruiksinstructies bevat voor gebruiksverantwoordelijken (Art. 13) - Effectief menselijk toezicht mogelijk is (Art. 14) - Het systeem passende niveaus van nauwkeurigheid, robuustheid en cybersecurity bereikt (Art. 15) Opzetten kwaliteitsmanagementsysteem volgens (Art. 17) Opstellen EU-conformiteitsverklaring en aanbrengen CE-markering volgens Art. 43–49 (Art. 16 lid g en h) Registreren in EU-databank volgens Art. 49 (Art. 16 lid i)	Verplichtingen van gebruiksverantwoordelijken van AI-systemen met een hoog risico (Art. 26): - Nemen van passende organisatorische en technische maatregelen om gebruik conform instructies te waarborgen (Art. 26 lid 1) - Opdragen/realiseren menselijk toezicht (Art. 26 lid 2) - Analyseren en controleren relevantie en representativiteit inputdata (voor zover die onder eigen controle vallen) (Art. 26 lid 4) - Monitoren en rapporteren van risico's en incidenten (Art. 26 lid 5) - Bewaren automatisch gegenereerde logs AI-toepassing (voor zover die onder eigen controle vallen) (Art. 26 lid 6) - Informeren werknemers bij gebruik AI-toepassing op de werkplek (Art. 26 lid 7) - AI-toepassing registreren in EU-databank (Art. 26 lid 8 / Art. 49) Beoordelen van de gevolgen voor grondrechten van AI-systemen met een hoog risico (FRIA) (Art. 27)
Beperkt risico	Ontwerp zodanig dat transparantie mogelijk is (Art. 50)	- Informatieplicht: kenbaar maken AI-gebruik (bijv. chatbot) (Art. 50 lid 1) - Labeling deepfakes/synthetische content (Art. 50 lid 2–4)
Algemene verplichtingen	- AI-geletterdheid (Art. 4)	- AI-geletterdheid (Art. 4)

**Let op: dit is een versimpelde juridische weergave van de verplichtingen uit de AI-verordening die niet volledig is

Stap 4: Aanvullende risicobeoordeling



WAAROM?

Voor hoog-risico AI-toepassingen is een risicobeoordeling verplicht, maar ook voor andere AI-toepassingen heeft het doen van een aparte risicobeoordeling meerwaarde. Met een risicobeoordeling kijk je naar specifieke context van je instelling en onderwijskundige toepassing. Lang niet alle mogelijke effecten en risico's worden door de wettelijke verplichtingen ondervangen. De AI-verordening is generiek, terwijl risico's sterk afhangen van de onderwijscontext, het concrete gebruik en de doelgroep.



MET WIE?

De aanvullende risicobeoordeling moet gezamenlijk door een multidisciplinair team worden gedaan :

- **De eerste lijn** moet de risico's in de praktijk beheersen en zorgen voor betrokkenheid door de risico-eigenaar. **De tweede lijn** is betrokken om kennis in te brengen welke risico's zich kunnen voordoen, vanuit verschillende disciplines zoals privacy, security en architectuur.
- **Mogelijke stakeholders zijn:** een manager, een ethisch adviseur, data-scientist of een HR-medewerker, en de groep die wordt blootgesteld aan de risico's, bijvoorbeeld in de vorm van studenten(vertegenwoordiging)



HOE?

De aanvullende risicobeoordeling bestaat, in lijn met de ISO 31000 voor risicomanagement en de [Toolkit risicobeoordeling van SURF](#), uit drie stappen:

1. **Identificeren** - vinden, herkennen en beschrijven van risico's.
2. **Analyseren** - inschatten van de waarschijnlijkheid en de ernst van de impact.
3. **Evalueren** - bepalen of en zo ja welke aanvullende acties of maatregelen vereist zijn.

Op de volgende pagina's staat een nadere toelichting en uitwerking van deze stappen.



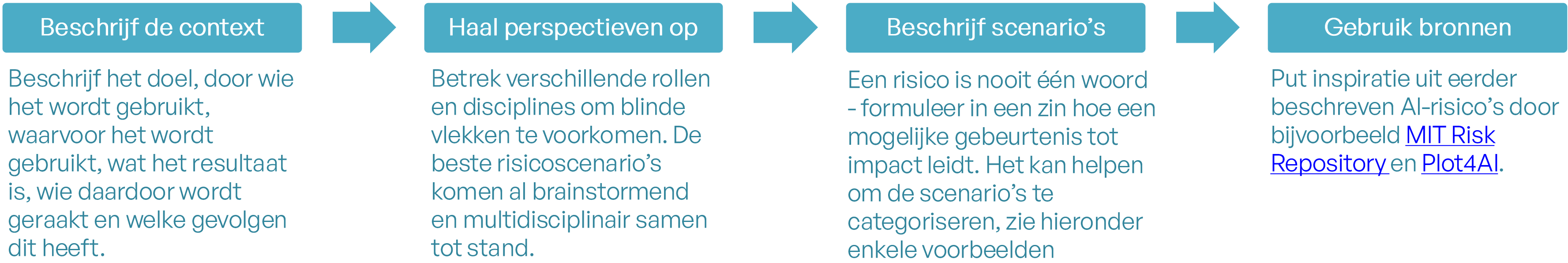
RESULTAAT:

De output van de aanvullende risicobeoordeling is een gedeeld beeld van de geïdentificeerde risico's, de waarschijnlijkheid van optreden en de mogelijke impact binnen de context van het gebruik. Op basis van de risico-evaluatie ontstaat inzicht in welke risico's acceptabel zijn en welke aanvullende beheersmaatregelen nodig zijn. Deze uitkomst vormt input voor de laatste stap: het implementeren, opvolgen van passende beheersmaatregelen en monitoring.

Aanvullende risicobeoordeling – (1) Risico-identificatie

Doel: In kaart brengen welke risico's de AI-toepassing kan hebben binnen de onderwijscontext. Dit resulteert in begrijpelijke en gedragen risico-scenario's die als basis dienen voor analyse, beoordeling en (later) besluitvorming over mogelijke maatregelen.

Hoe pak je het aan?



Mogelijke risicocategorieën

Bias

AI-systemen kunnen vooroordelen bevatten, wat kan leiden tot oneerlijke beoordelingen en ongelijke kansen.

Privacy en data

Studentgegevens, toetsresultaten of stage-informatie kunnen onbedoeld gedeeld of gelekt worden.

Transparantie

Onvoldoende inzicht in hoe AI tot beoordelingen komt, bemoeilijkt uitleg, controle en correctie.

Leerproces

Studenten worden afhankelijk van AI, waardoor het leerproces onder druk komt en vaardigheden afnemen.

Toetsing

Validiteit van toetsen en examineren komt onder druk te staan: meet de toets wat het moet meten?

Generieke risico's bij AI

Uniek voor het onderwijs

Aanvullende risicobeoordeling – (2) Risico-analyse

Doel: De risicoanalyse heeft als doel om de geïdentificeerde risico's nader te duiden door inzicht te geven in de waarschijnlijkheid dat een risico optreedt (kans) en de impact die dit heeft. Hiermee wordt zichtbaar welke risico's relatief zwaar wegen en dus prioriteit verdienen in de verdere besluitvorming.

Hoe pak je het aan?

Beoordeel kans

Het (kwalitatief) beoordelen van de waarschijnlijkheid van een risicoscenario kan door het stellen van de volgende vragen:

- Hoe waarschijnlijk is het dat het risico uitkomt?
- Zijn er eerdere incidenten of ervaringen die dit onderbouwen?
- Hoe effectief zijn de huidige beheersmaatregelen?



Beoordeel impact

Daarnaast is het belangrijk om de mogelijke impact in te schatten. Beantwoord bijvoorbeeld de volgende vragen:

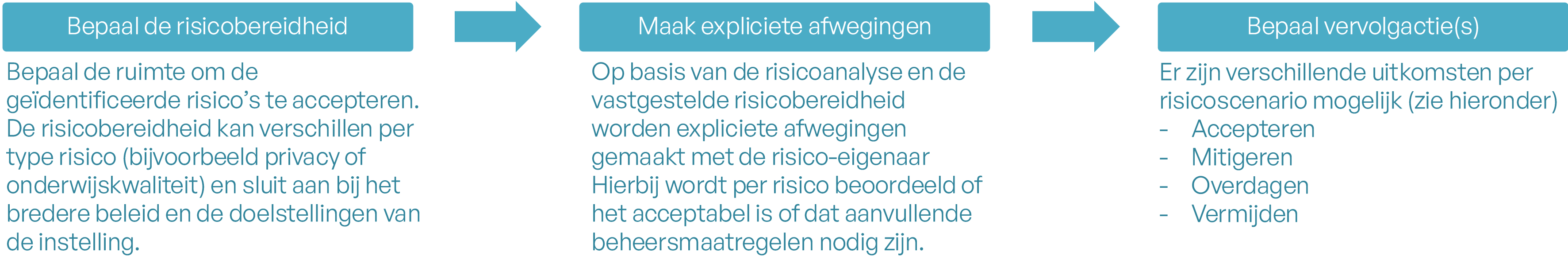
- Wat zijn de mogelijke gevolgen als het risico optreedt?
- Om hoeveel betrokkenen gaat het die geraakt worden?
- Wat is de ernst van de impact? (zie hieronder een illustratieve matrix die je kunt gebruiken)

Matrix die kan helpen bij het beoordelen van waarschijnlijkheid en impact			
Kans / Impact	Beperkt	Aanzienlijk	Kritiek
Onwaarschijnlijk			
Mogelijk			
Kansrijk			

Aanvullende risicobeoordeling – (3) Risico-evaluatie

Doel: Nadat je de risicoscenario's hebt beschreven en de waarschijnlijkheid en impact hebt geanalyseerd, kan je evalueren welke risico's je bereid bent om te nemen en welke beheersmaatregelen je wilt treffen. Dit is de laatste stap van de risicobeoordeling, namelijk de risico-evaluatie. Het is belangrijk om hierbij altijd de risico-eigenaar te betrekken.

Hoe pak je het aan?



Mogelijke vervolgactie(s)

Accepteren

Geen aanvullende maatregelen nodig, wel monitoren. Indien context verandert, risicobeoordeling opnieuw uitvoeren.

Mitigeren

Neem aanvullende beheersmaatregelen om de kans en/of de impact van het risico te verkleinen.

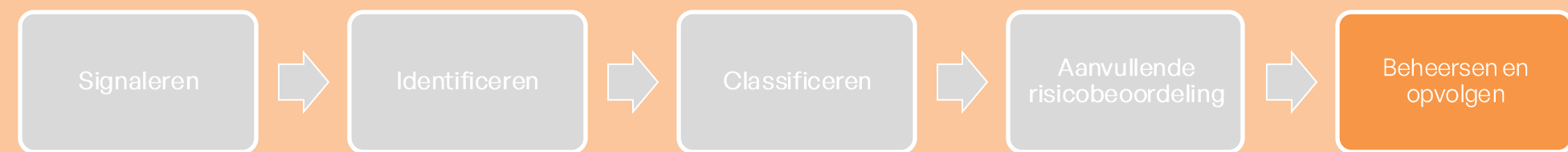
Overdragen

Draag het risico over aan een derde partij, bijvoorbeeld door te verzekeren of uit te besteden.

Vermijden

Het risico is niet acceptabel of onvoldoende beheersbaar. Neem het AI-systeem niet in gebruik, of faseer het uit (indien al in gebruik).

Stap 5: Beheersen en opvolgen



WAAROM?

Het doel van deze stap is om ervoor te zorgen dat de maatregelen voor de geïdentificeerde risico's ook daadwerkelijk worden uitgevoerd, gemonitord en bijgestuurd. Hetzelfde geldt voor de wettelijke verplichtingen; deze zijn door middel van dit raamwerk bepaald, maar er moet ook actie worden genomen om aan de verplichtingen te voldoen. Dit vraagt om beheersing en opvolging. Hierbij is het van belang dat een duidelijke taakverdeling wordt gemaakt en duidelijk is wie waarvoor verantwoordelijk is.

MET WIE?

- **De eerste lijn** is betrokken om een keuze te maken over het in gebruik nemen van het AI-systeem en de manier waarop risico's in de praktijk worden beheerst. **De tweede lijn** is betrokken voor compliance-advies. Alle voorgaande stappen hebben input opgeleverd over de verplichtingen en risico's en deze moeten beheerst worden.
- **Mogelijke stakeholders zijn:** een inkoop- contactadviseur (als het systeem wordt ingekocht), een leverancier/ontwikkelaar, een implementatie-adviseur of ondersteunend personeel.

HOE?

De uitvoering is vormvrij - iedere onderwijsinstelling heeft eigen afspraken, processen en systemen voor risicobeheersing en risicomanagement. Het ligt voor de hand om AI-risico's hierin te integreren en geen aparte systemen in te richten specifiek voor AI-toepassingen.

Een mogelijke manier om risico's te registreren is door het vullen van een algoritmeregister. Werkgroep Algoritmeregister heeft hiervoor een template opgesteld. Of je maakt gebruik van de use-case-card-methode. Een voorbeeld hiervoor staat op de [volgende pagina](#).

RESULTAAT:

Stap 5 'beheersen en opvolgen' vormt de laatste stap in het risicobeoordelingsraamwerk. Na deze stap:

- zijn de risico's vastgelegd
- is duidelijk wie waarvoor verantwoordelijk is
- worden risico's actief beheerst met passende maatregelen
- wordt een cyclisch proces van monitoring en bijsturing gevolgd op basis van het risicomanagementproces van de instelling

Gebruik van use-case-cards voor vastlegging en opvolging

Use case cards

Nu je alle stappen in het raamwerk hebt doorlopen, is het belangrijk om informatie over de AI-toepassing en risico's vast te leggen, zodat je deze straks goed kunt beheersen. Hiervoor kan je gebruik maken van de use-case-card-methode. Hieronder staat een voorbeeld van een use case card die je na het volgen van dit raamwerk kunt invullen; gebruik de laatste kolom in en volg de schrijfinstructies die tussen [haken] staat. De use case card is gebaseerd op het werk van ICTRecht, de werkgroep beleidspiramide mbo en het Privacy Expertise Centrum van SURF.

Stap raamwerk	[Titel van de toepassing waarvoor je de use case card invult]	
1. Signaleren	Context van gebruik	[Beschrijf het proces of systeem waar de AI-toepassing deel van gaat uitmaken]
	Scope	[Beschrijf de rol van de AI-toepassing en de grenzen van het gebruik]
	Beleidsdoeleinden	[Beschrijf welke beleidsdoeleinden de organisatie voor ogen heeft en hoe AI die doeleinden gaat vervullen]
2. Identificatie	Definitie	[Valt de toepassing wel of niet onder de definitie van een AI-systeem volgens de AI-verordening]
3. Classificatie	Ingenomen rol	[Aanbieder of gebruiksverantwoordelijke]
	Risiconiveau	[Beschrijf onder welk risiconiveau de AI-toepassing valt en/of er sprake is van uitzonderingsgronden]
	Wettelijke verplichtingen	[Beschrijf de wettelijke verplichtingen uit de AI-verordening waar de toepassing uitgaande van het risiconiveau aan dient te voldoen en hoe je daar uitvoering aan geeft]
4. Aanvullende risicobeoordeling	Risico-identificatie	[Beschrijf de risico-scenario's van de AI-toepassing]
	Risico-analyse	[Beschrijf de kans en impact van de geïdentificeerde risico's van de AI-toepassing]
	Risico-evaluatie	[Beschrijf welke risico's je bereid bent om te nemen, de maatregelen die je treft en wie deze uitvoert]
5. Beheersing & opvolging	Rollen en verantwoordelijkheid	Eerste lijn [Beschrijf wie de eindgebruiker(s) is, de risico-eigenaar en de eindverantwoordelijke eigenaar]
		Tweede lijn [Beschrijf welke rollen vanuit de tweede lijn ondersteunen en adviseren]
	Status raamwerk	[Vermeld welke stappen zijn afgerond, overgeslagen (waarom?) of nog openstaan]
	Datum	[Datum waarop waarop raamwerk doorlopen is]

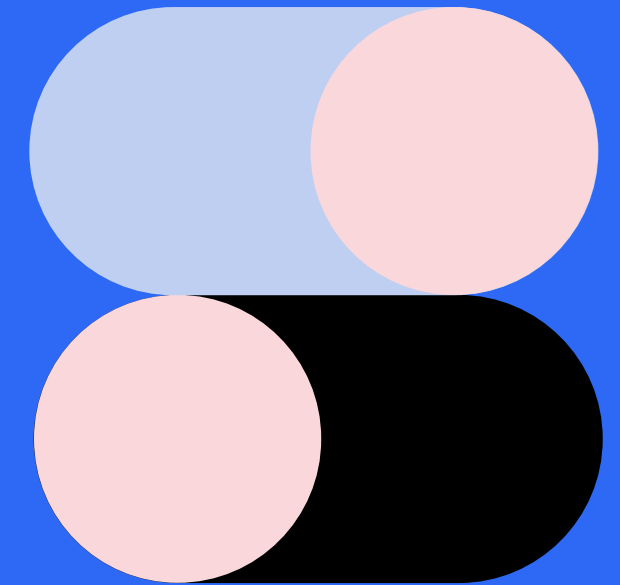
Risicobeoordeling als continu proces

Nu je het raamwerk hebt doorlopen, heb je inzicht in het risiconiveau van het AI-systemen dat je gebruikt en de bijbehorende verplichtingen uit de AI-verordening. Ook heb je tools en methodes aangereikt gekregen die kunnen helpen met het inschatten en beoordelen van risico's. Het raamwerk is nu eenmalig doorlopen, maar risicobeoordeling is een continu proces. Het raamwerk kan keer op keer gebruikt worden, voor nieuw gesignaleerde AI-systemen binnen je instelling, maar ook om de risico's van de al gebruikte AI-systemen opnieuw te beoordelen en te actualiseren.

Ook het raamwerk zelf wordt doorontwikkeld, op basis van nieuwe inzichten of praktijkervaringen. Ook al zijn alle stappen doorlopen, risicobeoordeling is eigenlijk nooit echt klaar. We moedigen je aan om hier binnen je instelling continu mee bezig te blijven en dit raamwerk er later nog eens bij te pakken om weer tot nieuwe inzichten te komen.

**De eerder in dit raamwerk genoemde hulpmiddelen dienen uitsluitend ter illustratie en vormen geen volledige lijst. Er zijn mogelijk andere geschikte hulpmiddelen beschikbaar of deze komen op termijn beschikbaar. Organisaties kunnen gebruikmaken van alternatieve hulpmiddelen die aansluiten bij hun eigen context en behoeften.*

Contact



- Dit document wordt beheerd door Npuls. Voor nadere informatie neem contact op via aiendata@npuls.nl o.v.v. RAISE.
- Het raamwerk wordt actief beheerd en regelmatig bijgewerkt om de meest actuele informatie te bieden. Je vindt de meest recente versie hier: <https://edu.nl/6huua>

Colofon

Dit raamwerk is ontwikkeld door Npuls in samenwerking met de werkgroep RAISE, en getoetst bij een groep stakeholders. De ontwikkeling van het raamwerk is begeleid door onderzoeks- en adviesbureau PBLQ. Op alfabetische volgorde:

Auteurs / werkgroep RAISE:

- Nicole van Deursen (SURF)
- Henriette Dietz (Erasmus Universiteit Rotterdam)
- Irene Eegdeman (ROC van Amsterdam-Flevoland)
- Matijs van Hilst (MBO Digitaal)
- Arvin Khozoei (Technische Universiteit Delft)
- Yvette Roman (Universiteit Utrecht)
- Saskia Stout (Hogeschool Windesheim)

Reviewers / stakeholders:

- Jack Le Doux (Vista college, MBO Digitaal)
- Martijn Hendriks (Hogeschool Arnhem Nijmegen)
- Janneke Kuipers (Tilburg University)
- Daisy Oolbekkink (Universiteit Twente)
- Dylan Scheurs (ROC Midden Nederland)
- Eva Willems (Haagse Hogeschool)

Begeleiding:

- Duuk Baten (Npuls)
- Tess van Delft (PBLQ)
- Marian Biekart (PBLQ)
- Juliette Mies (PBLQ)
- Stephan Okhuijsen (Npuls)
- Jinne Samsom (PBLQ)
- Linda Torn (PBLQ)