

AI als spiegel en motor – besturen van binnenuit in een algoritmisch tijdperk

René de Baaij, Groesbeek, 19 januari 2025

Leeswijzer

- **Voor wie:** RvB (Raad van Bestuur) / RvC (Raad van Commissarissen), MT (managementteam) en publieke bestuurders; focus op strategie, legitimiteit en zorgplichten.
- **Hoe te lezen:** secties 1–7 bouwen de gedachtegang op (mens & systeem); secties 9–11 vormen het instrumentarium. Lees Proloog (sectie 0) voor context; spring voor implementatie gerust naar 9–11 en blader terug voor verdieping.
- **Wat u meeneemt:** drie beslislijnen (waarde, risico, legitimiteit), één ritme (monitor–challenge–audit) en vijf artefacten (modelkaart, logboek, risicodossier, gebruikershandleiding, exit-plan).

Proloog – De kamer, het scherm, de blik

Het is vroeg. Het scherm licht op nog vóór de vergaderzaal zich vult. De eerste vraag van de dag is geen begroeting maar een prompt. Er rolt een antwoord uit dat beter samenhangt dan de notities van gisteren. Iemand zucht van opluchting, iemand anders fronst, een derde voelt onverklaarbare irritatie. In die drie microreacties ligt het speelveld bloot: AI als belofte, als bedreiging, als spiegel. Wat hier gebeurt is niet louter technisch of procedureel; het is ook psychologisch, relationeel, bestuurlijk en juridisch. Het raakt hoe u richting geeft, hoe teams betekenis maken en hoe de organisatie verantwoordelijkheid draagt.

Deze longread verbindt twee lenzen die zelden in één verhaal samenvallen: het psychodynamisch, relationeel en humanistisch fundament én het bestuurlijk, bedrijfskundig en juridisch perspectief. We kijken tegelijk naar binnen en naar buiten, naar onderstroom en bovenstroom. Naar wat AI met mensen doet, en wat mensen met AI doen; naar stuur, structuur en verantwoording.

De kernboodschap is eenvoudig en scherp: **AI verandert de mens niet van soort, maar vergroot wat er al is – en maakt de gevolgen bestuurlijk en juridisch expliciet.**

Kernstatements (met korte toelichting)

- **AI vergroot wat er al is.**
Het versterkt patronen, versnelt ritmes en maakt impliciete normen expliciet in code. Dat levert scherpte op, maar ook vergrote blinde vlekken als tegenkracht ontbreekt.
- **Neutraliteit is schijn zonder governance.**
Een model is geen arbiter maar een ontwerpkeuze met aannames, data en drempels.

Leg herroepbaarheid en betwistbaarheid vast, anders wordt 'efficiëntie' een alibi.

- **Relatie gaat vóór techniek.**

Wat het systeem met mensen doet en mensen met het systeem vormt de cultuur.

Dat vraagt om dialoog, uitzondering en herstel als structurele functies.

- **Legitimiteit is bewijsbaar handelen.**

Zorgvuldigheid telt pas als ze navolgbaar is in logboeken, kaarten en besluiten.

Uitleg en herstel zijn geen bijzaak maar onderdeel van het product.

1. Binnenwereld en bovenstroom – projecties op het algoritme

In klassieke verhoudingen werden almacht, zekerheid en alwetendheid op de leider geprojecteerd. Nu schuift het algoritme in die rol. Het verschijnt als het nieuwe "object": slim, snel, schijnbaar neutraal. Die schijn van neutraliteit maakt het aantrekkelijk als drager van idealisatie en ontkenning. "Het model zegt het" is uitgegroeid tot een modern verdedigingsmechanisme – een manier om schaamte, twijfel of morele last te ontwijken.

Psychodynamisch gezien is dat geen afwijking maar voorspelbaar: overdracht en contra-overdracht spelen zich net zo goed af **tussen mensen en systemen** als tussen mensen onderling. We praten met het systeem, maar eigenlijk met onze behoefte aan zekerheid. We luisteren naar de score, maar eigenlijk naar onze angst om tekort te schieten.

Bestuurlijk werkt dit door. Waardecreatie verschuift van individuele expertise naar de **configuratie** van data, processen, modellen en governance. Juridisch schuift de vraag mee: niet "wie heeft het gedaan?", maar "hoe is het systeem ingericht, wie is eigenaar, welke zorgplichten golden en is herstel mogelijk?"

In termen van motivatie raakt dit autonomie, competentie en verbondenheid (Self-Determination Theory; zie noot 1).

Drie eenvoudiger, daarom moeilijkere vragen horen vanaf nu bij ieder AI-dossier:

1. **Wat doet dit systeem met onze rollen en relaties?** Wie krijgt stem, wie wordt stil?
2. **Waar gebruiken we data en modellen om contact te vermijden in plaats van te verdiepen?**
3. **Hoe blijven we geworteld in empathie en proportionaliteit juist wanneer beslissingen worden 'geoptimaliseerd'?**

Wie ze ontwijkt krijgt snelheid zonder richting en compliance zonder legitimiteit.

2. AI als systeemknoop – patronen, macht en variëteit

AI is geen los gereedschap; het is een **systeemknoop** die nieuwe feedbacklussen, afhankelijkheden en machtsstructuren creëert. Vanuit systeemdenken en complexiteitsleer verschuift de kernvraag: niet of AI “werkt”, maar **welke patronen het versterkt of onderdrukt**.

Requisite variety wordt concreet (zie noot 8). Organisaties die complexe realiteit willen hanteren, hebben variatie in perspectieven nodig. Veel modellen reduceren juist: ze scoren, sorteren, normaliseren. Dat is efficiënt, maar ook een bron van **cultural drift** naar conformiteit en risicomijding. Minder afwijking, minder tegenspraak, minder creativiteit – met als bestuurlijk neveneffect: blind spots die pas zichtbaar worden in incidenten en klachten.

In systemische taal **codeert AI impliciete normen**. Macht en habitus verschuiven naar datasets, features en thresholds. Wie niet oplet, bestendigt uitsluiting als “objectieve logica”. Daarom is governance niet primair een document, maar een **relationeel ontwerp**: wie definieert de data, wie kiest de modellen, wie mag uitzonderen en wie mag overrulen? Dit perspectief is sociomaterieel van aard (zie noot 6) en sluit aan bij complex responsive processes (zie noot 7).

Voor u als bestuurder volgt hieruit een simpele regel met harde consequenties: **AI-trajecten zijn verandertrajecten**. Het zijn ingrepen in rollen, processen, identiteit en verantwoording. Wie ze als IT (informatietechnologie)-project uitvoert, krijgt technische opleveringen en sociale nevenschade.

3. Strategie en operatie – snelheid, betekenis en meetlatten

AI kan productiviteit verhogen, doorlooptijden verkorten en kwaliteit voorspelbaarder maken. Maar zonder **doelarchitectuur** vervalt het in losse pilots en demowinst. Strategisch is de volgorde: **probleem → waarde → risico → ontwerp**. Eerst helder definiëren welk bedrijfsprobleem en welke waardeformule, dan pas tools. Anders anders wordt technologie oplossing en mens probleem.

Operationeel vraagt AI om een zichtbaar **ketenontwerp**: van vraag tot effect. Intake, ontwerp, validatie, go-live, monitoring, hertraining en uitfasering vormen één ritme. Drempelwaarden en retrain-criteria worden vooraf vastgelegd. Prompts, features, hyperparameters en versies worden

assets met eigenaarschap. U ziet: dit is geen hobbyhoek, dit is service management voor besluitvorming.

Leren is de zwakke plek. AI excelleert in **single-loop**: beter worden in wat we al doen. **Double-loop** – de vraag of we de juiste dingen meten, de passende succescriteria hanteren – blijft een menselijk en organisatorisch proces. Wie die laag schraapt, vergroot blindheid met hoge precisie. Het gevolg: mooie dashboards, verkeerde beslissingen (Kolb; zie noot 3; Argyris & Schön; zie noot 4; Mezirow; zie noot 5).

Daarom organiseren lerende organisaties drie praktijken: (1) **reflectie op uitkomsten** – niet alleen klopt het, maar wat betekent dit?; (2) **dialogische besluitvorming** – data als gespreksstarter, niet als arbiter; (3) **co-evolutie van vaardigheden** – digitale geletterdheid plus oordeelsvermogen, morele sensitiviteit en juridische alertheid.

4. Rechtmatigheid en legitimiteit – de normatieve infrastructuur

AI werkt binnen een **normatief landschap** dat in Europa gelaagd en in beweging is. De basisprincipes van gegevensbescherming staan al jaren: rechtmatigheid, doelbinding, dataminimalisatie, juistheid, opslagbeperking, integriteit en verantwoordingsplicht.

Daaromheen liggen regels voor platforms, data-toegang en -portabiliteit, veiligheid en zorgplichten in essentiële sectoren (zie noten 9–13). Specifieke AI-kaders vragen beheerste risicoanalyse, documentatie, monitoring en correctiemogelijkheden.

Het bestuurlijke werkwoord is **bewijsbaar**. Niet alleen zorgvuldig zijn, maar kunnen laten zien dat u zorgvuldig wás: modelkaarten, logboeken, versiebeheer, impactassessments, override-besluiten met motivering, incident-notities met herstel. Leg tegenkracht vast: wie **mag** het systeem stoppen, en **wanneer**?

Legitimiteit is meer dan legaliteit. Ze rust op **uitlegbaarheid, betwistbaarheid en herstel**. Wie automatiseert moet ook zichtbaar maken **wat níet wordt geautomatiseerd**. Trek rode lijnen: geen gedragssturing zonder noodzaak en proportionaliteit; geen black-box in contexten met hoge foutkosten en lage uitlegbaarheid; geen dataverzameling zonder noodzaak en heldere grondslag. Maak dit expliciet, in taal die klanten, cliënten, inwoners en medewerkers begrijpen.

5. Identiteit en belichaamd werken – spanning kunnen dragen

AI raakt professionele identiteit. Als systemen analyseren, adviseren en creëren, wat betekent het om expert of leider te zijn? De opgave verschuift van **meer weten** naar **meer kunnen dragen**: spanning, onzekerheid, traagheid, conflict en morele ambiguïteit.

Spirituele en humanistische taal krijgt hier een concrete functie: richting houden in een wereld die mechanisch maakbaar lijkt. De kernvragen worden: **wat doen we níet, ook al kan het?** Welke toepassingen weigeren we omdat ze relaties, waardigheid of toekomst uithollen?

Belichaamd werken is randvoorwaarde. AI versnelt ritme en verhoogt informatiedruk. Het lichaam betaalt als eerste. Teams die ritme, stilte, herstel en reflectie organiseren, houden koers. Teams die die laag overslaan, branden op – en dat is niet alleen een welzijnskwestie, maar ook een bestuurlijk risico: uitval, fouten, afbrokkelend vertrouwen.

6. Methodische verankering – onderzoek als bewustzijnspraktijk

De meeste AI-gesprekken stranden in toolkeuze, pilots en risico-lijsten. Wat ontbreekt is **methodische verankering**: een manier van werken waarin reflectie, dialoog en betekenisgeving net zo serieus worden genomen als accuraatheid en uptime.

Kwalitatieve methoden – interview, casuïstiek, discours- en narratieve analyse – zijn onmisbaar om te begrijpen wat AI doet met ervaring en interactie. Cijfers tonen patronen, verhalen tonen betekenis. Combineer dit met **action research**: hypothesen formuleren, effecten volgen, aanpassingen doorvoeren, herbeslissen. Elk AI-traject wordt zo een leerproces.

Aan de formele kant vraagt dit om een **AI-managementsysteem (zie noten 14–19)**: beleid en standaarden, processen voor risico, validatie, monitoring, incidentrespons, uitfasering; competenties (model risk, data-ethiek, juridische duiding); assurance (interne audit, externe toetsing, maturity-meting). En – cruciaal – **due diligence** richting leveranciers: databronnen en licenties, evaluatiesets en bias, red-teaming, content-authenticiteit, portabiliteit en exit-opties. Contracteer auditrechten en sancties bij non-compliance.

Methodologie wordt zo ook **ethiek**: hoe komen we tot ons oordeel, wie wordt gehoord, welke data

tellen mee en welke niet? Het eigenaarschap over betekenis moet bij mensen blijven, niet bij de leverancier. Dit sluit aan bij internationale principes en richtsnoeren (zie noten 20–21).

7. Paradoxen en condities – het noodzakelijke ongemak

De bestuurlijke werkelijkheid is een veld van spanningen. Al maakt ze zichtbaar en onvermijdelijk:

- **Efficiëntie ↔ veerkracht.** Sneller en goedkoper kan fragieler zijn. Conditie: kies veerkracht waar verstoring kostbaar is; kies efficiëntie waar herstel makkelijk is.
- **Centralisatie ↔ autonomie.** Eén platform geeft schaal, maar vermindert lokale wijsheid. Conditie: centrale standaarden, lokale variatie in toepassing.
- **Automatisering ↔ menselijke waardigheid.** Delegatie ontlast, maar kan ontmenselijken. Conditie: duidelijke grenzen voor delegatie, met betwistbaarheid en herstel.
- **Dataminimalisatie ↔ personalisatie.** Meer data is niet altijd beter. Conditie: doelbinding expliciet, databehoud beperkt, effect aantoonbaar.
- **Uitlegbaarheid ↔ performance.** Complexe modellen presteren soms beter, maar zijn minder navolgbaar. Conditie: match modelkeuze aan foutkosten en context.
- **Openheid ↔ veiligheid.** Transparantie vergroot vertrouwen, maar ook aanvalsvlak. Conditie: differentieer publiek, intern en vertrouwelijk.

Spanningen verdwijnen niet door beleid; ze vragen ritme en gesprek. Besturen is **koers houden temidden van tegenstrijdige waarheden**.

9. Minimal viable governance – klein aantal dingen heel goed

Grote ambities slagen door **weinig dingen consequent te doen**. Een werkend minimum waar u morgen mee kunt beginnen:

Rollen

- **Systeemeigenaar (accountable)** – mandaat, middelen, verslag.
Houdt het mandaat vast en draagt integraal verantwoordelijkheid.
Borgt besluitlijnen, middelen en verslaglegging richting bestuur.
- **Model owner** – prestaties, bias, lifecycle.
Stuurt prestatie, bias en veroudering van het model end-to-end.
Bewaakt validatie, monitoring en tijdige retraining of retirement.

- **Data-eigenaar** – definities, kwaliteit, toegang.
Definieert bronnen, kwaliteit en toegangsrechten van data.
Bewaakt doelbinding, minimalisatie en datakarakteristieken voor modellen.
- **Risk & compliance** – kaders, toetsen, escalaties.
Vertaalt wet- en normenkaders naar werkbare eisen in het proces.
Orkestreert toetsen, uitzonderingen en escalaties met tijdigheid.
- **Security** – dreigingen, weerbaarheid, incidenten.
Beoordeelt dreigingen en kwetsbaarheden over de volledige keten.
Borgt detectie, respons en herstel inclusief leveranciersafhankelijkheden.
- **Business owner** – doel, waarde, effecten in de praktijk.
Verankert probleemdefinitie, waarde en effect in de operatie.
Organiseert feedback uit de praktijk en besluit over stop/go.

Ritme

- **Maandelijks:** modelmonitoring en incidentreview.
Beoordeelt performance, drift en incidenten met concrete acties.
Normaliseert kleine correcties zodat grote fouten uitblijven.
- **Elk kwartaal:** challenge-sessie met onafhankelijke tegenkracht.
Biedt georganiseerde tegenspraak en alternatieve aannames.
Voorkomt groepsdenken en verfijnt drempelwaarden en scope.
- **Halfjaarlijks:** audit op documentatie, drift, fairness en rechten.
Geeft assurance op documentatie, fairness en rechtsbescherming.
Maakt verbeterbesluiten navolgbaar voor audit en stakeholders.
- **Bij iedere wijziging:** impactassessment, testrapport, go/no-go door board met override-criteria.
Voert vooraf impactanalyse en tests uit met duidelijke criteria.
Legt motivering van go/no-go en eventuele overrides traceerbaar vast.

Artefacten

- **Modelkaart** (doel, data, grenzen, performance, fairness, uitlegbaarheid).
Beschrijft doel, grenzen, data, prestaties en aannames.
Vormt basis voor uitleg, audit en verantwoord gebruik.
- **Logboek** (versies, parameters, overrides, incidenten, herstel).
Houdt alle versies, parameters, overrides en incidenten bij.
Maakt terugkijken, leren en aansprakelijkheid mogelijk.
- **Risicodossier** (foutkosten, context, mitigaties, residual risk).
Vertaalt foutkosten en context naar mitigaties en residual risk.
Helpt bij besluitvorming over modelkeuze en delegatiegraad.

- **Gebruikershandleiding** in begrijpelijke taal (rechten, bezwaar, contact).
Legt rechten, bezwaar en contact voor gebruikers helder uit.
Verhoogt acceptatie en maakt betwistbaarheid praktisch en laagdrempelig.
- **Exit-plan** (vendor lock-in, portabiliteit, escrow).
Definieert portabiliteit, escrow en migratiestappen bij beëindiging.
Verkleint vendor lock-in en continuïteitsrisico's.

Dit alles is compact te houden. Het is geen papierfabriek; het is **traceerbaar vakmanschap**.

10. Van principe naar praktijk – drie ontwerp vragen

A. Waar plaatsen we de menselijke tegenkracht?

In de loop (realtime), on the loop (periodiek toezicht) of above the loop (kaderstellend). Kies bewust per use case en leg uit waarom.

Aan-de-lijn verlaagt u foutkosten bij hoge impact en lage tolerantie.

Boven- en on-the-loop borgen proportionaliteit en leerbaarheid over tijd.

B. Hoe wordt uitzondering normaal?

Elke regel kent grensgevallen. Maak uitzonderen legitiem: wie mag het, hoe motiveert men, hoe wordt ervan geleerd?

Leg procedure, criteria en motiveringsplicht vast – geen coulance per geval.

Documenteer, deel patronen en voer periodiek verbeterbesluiten door.

C. Welke maat van variëteit is nodig?

Standaardiseer waar veiligheid en schaalbaarheid dat vragen. Laat variatie waar contextkennis kwaliteit toevoegt. Gebruik governance om **verschil mogelijk en navolgbaar** te maken, niet om het uit te gummen.

Te veel standaard kost kwaliteit; te veel variatie kost schaal.

Definieer keuzeruimte en afwijkingroutes expliciet per context.

11. Wat nu, wat later, wat niet – handelingsperspectief

Nu (0–3 maanden)

1. **Inventariseer** alle AI-toepassingen en modellen; wijs eigenaren aan.
Zonder zicht op de feitelijke footprint blijft governance theoretisch en risico ongericht.

Een sobere inventaris voorkomt schaduw-IT (informatietechnologie), versnelt audits en maakt prioriteren mogelijk.

2. **Definieer risk appetite** en kritieke kwaliteitsattributen (veiligheid, fairness, privacy, uitlegbaarheid).

Zonder expliciete grenzen worden compromissen impliciet en willekeurig.

Heldere drempels sturen ontwerpkeuzes en versnellen besluitvorming bij incidenten.

3. **Stel stopcriteria en override-bevoegdheid** vast; organiseer een praktische kill-switch.

Een systeem dat niet veilig kan worden gestopt is niet veilig om te starten.

Duidelijke bevoegdheden beperken schade, versnellen herstel en verhogen vertrouwen.

4. **Start impactassessments** voor de top drie use-cases; leg datastromen, aannames en foutkosten vast.

Impact gaat over mensen, processen en rechten—niet alleen over metrics.

Vroeg inzicht voorkomt dure herontwerpen en faciliteert uitlegbaarheid naar buiten.

5. **Regel het ritme:** maandelijkse monitoring en kwartaal-challenge met onafhankelijke tegenkracht.

Zonder ritme verschuift aandacht naar incidenten; met ritme groeit lerend vermogen.

Onafhankelijke tegenspraak voorkomt groepsdenken en normaliseert het corrigeren van koers.

Later (6–12 maanden)

1. **Bouw het AI-managementsysteem:** processen, rollen, audit-trails, maturity-meting.

Kwaliteit ontstaat niet uit intentie maar uit herhaalbare werkpraktijken.

Een licht, traceerbaar systeem maakt prestaties schaalbaar zonder bureaucratie.

2. **Vendor due diligence:** licenties, datasets, evaluaties, red-teaming, content-authenticiteit, exit-opties.

Wat u uitbesteedt, blijft uw verantwoordelijkheid richting klant en toezichthouder.

Contractuele auditrechten en exit-mogelijkheden beperken lock-in en juridische kwetsbaarheid.

3. **Train leidinggevenden en sleutelfiguren** in oordeelsvorming onder onzekerheid en juridische duiding.

Technische kennis zonder oordeelsvorming vergroot precisie van verkeerde beslissingen.

Gezamenlijk taal en criteria ontwikkelen verhoogt consistentie en uitlegbaarheid.

4. **Integreer reflectie in de operatie:** casuïstiekbesprekingen, frontlinie-narratieven, verbeterbesluiten met motivering.

Verhalen geven betekenis aan data en maken effecten voelbaar en corrigeerbaar.

Door besluitmotieven vast te leggen, ontstaat een reproduceerbare kwaliteitscultuur.

Niet (tenzij uitzonderlijk onderbouwd)

- Onomkeerbare delegatie van hoog-impactbesluiten zonder menselijke tegenkracht.

Delegatie zonder rem tast waardigheid en rechtsbescherming aan bij fouten.

Beperk daarom de scope en borg betwistbaarheid en herstel vooraf.

- Black-box in contexten met hoge foutkosten en lage uitlegbaarheid.

Waar de prijs van een fout hoog is, is navolgbaarheid onderdeel van veiligheid.

Kies eenvoud of hybride modellen als uitleg essentieel is voor vertrouwen.

- Dataverzameling zonder noodzaak, grondslag en duidelijke bewaartermijn.

Meer data is niet vanzelf beter; het vergroot risico's zonder proportionele waarde.

Formuleer doelbinding, minimaliseer en stel tijdige verwijdering zeker.

- KPI's (kritieke prestatie-indicatoren) die perfect meten wat onbelangrijk is en niets zeggen over betekenis en legitimiteit.

Precisie zonder relevantie stuurt op schijnprestaties en ondermijnt vertrouwen.

Koppel meetpunten aan waarde, rechten en effecten in de leefwereld.

12. Slot – Besturen als kunst van het dragen

AI is geen buitenstaander die even binnenkomt; het is een **verlengstuk van onze vermogens en onze blinde vlekken**. Het vergroot ons kunnen en ons tekort. Besturen in dit tijdperk is de kunst van het **dragen**: spanning dragen, pluraliteit dragen, de traagheid die kwaliteit nodig heeft dragen. Het vraagt om verbeelding én discipline, om innerlijke oriëntatie én navolgbare processen.

De uitnodiging is eenvoudig en veeleisend:

Maak techniek tot bondgenoot van menselijkheid en rechtsstatelijkheid.

Ontwerp tegenkracht die werkt wanneer het erop aankomt.

Houd het gesprek gaande – juist waar het schuurt.

Als u die drie lijnen vasthoudt, wordt AI geen zwarte doos maar een heldere spiegel. Geen machinerie die mensen vervangt, maar een infrastructuur die professioneel oordeel, relationele kwaliteit en publiek vertrouwen **vergroot**. Dat is transformatie van binnenuit: stap voor stap, zichtbaar verantwoord, met een ritme dat vol te houden is.

Reflectieve vraag om mee af te sluiten:

Welke beslissing in uw organisatie kan morgen beter worden door minder automatisering en meer uitleg – en wie nodigt u vandaag nog uit om daar samen naar te kijken?

Notenapparaat & literatuur (selectie)

Theorie & begrippen (mens, leren, systeem)

1. Ryan, R. M., & Deci, E. L. (2000). *Self-Determination Theory and the Facilitation of Intrinsic Motivation, Social Development, and Well-Being*. **American Psychologist**, **55**(1), 68–78. https://selfdeterminationtheory.org/SDT/documents/2000_RyanDeci_SDT.pdf
2. Maslach, C., Schaufeli, W. B., & Leiter, M. P. (2001). *Job Burnout*. **Annual Review of Psychology**, **52**, 397–422. https://dspace.library.uu.nl/bitstream/handle/1874/13606/maslach_01_jobburnout.pdf
3. Kolb, D. A. (1984). **Experiential Learning: Experience as the Source of Learning and Development**. Prentice-Hall. (Overzicht: https://books.google.com/books/about/Experiential_Learning.html?id=zXruAAAAMAAJ)
4. Argyris, C., & Schön, D. A. (1978). *Organizational Learning: A Theory of Action Perspective*. Addison-Wesley. (Openbare versie: <https://archive.org/details/organizationalle00chri>)
5. Mezirow, J. (1991). *Transformative Dimensions of Adult Learning*. Jossey-Bass. (Samenvatting: <https://eric.ed.gov/?id=ED353469>)
6. Orlikowski, W. J. (2007). *Sociomaterial Practices: Exploring Technology at Work*. **Organization Studies**, **28**(9), 1435–1448. (Open access variant: <https://www.dhi.ac.uk/san/waysofbeing/data/data-crone-orlikowski-2007.pdf>)
7. Stacey, R. D. (2001). **Complex Responsive Processes in Organizations: Learning and Knowledge Creation**. Routledge. (Overzicht: <https://www.routledge.com/Complex-Responsive-Processes-in-Organizations-Learning-and-Knowledge-Creation/Stacey/p/book/9780415249195>)
8. Ashby, W. R. (1956). *An Introduction to Cybernetics*. Chapman & Hall. (Open access: <https://archive.org/download/AnIntroductionToCybernetics/AnIntroductionToCybernetics.pdf>)

Normenkader EU (peildatum 19 januari 2026, Europe/Amsterdam)

1. **AVG / GDPR** — Verordening (EU) 2016/679. Officiële publicatie: EUR-Lex. <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>
2. **DSA** — Verordening (EU) 2022/2065 (Digital Services Act). Officiële publicatie: EUR-Lex. <https://eur-lex.europa.eu/eli/reg/2022/2065/oj/eng>
3. **Data Act** — Verordening (EU) 2023/2854. Officiële publicatie: EUR-Lex. <https://eur-lex.europa.eu/eli/reg/2023/2854/oj/eng>
4. **NIS2** — Richtlijn (EU) 2022/2555. Officiële publicatie: EUR-Lex. <https://eur-lex.europa.eu/eli/dir/2022/2555/oj/eng>
5. **AI Act (EU)** — Implementatietijdlijn en toepassingsdata. European Parliament Research Service (At-a-Glance, 10 juni 2025; met overzichtsdatum juli 2024). <https://www.europarl.europa.eu/>

thinktank/en/document/EPRS_ATA\2025\772906

– Samenvattend: inwerkingtreding 1 augustus 2024; verbodsbepalingen van toepassing per 2 februari 2025; brede toepassingsdatum 2 augustus 2026 (bron: EPRS en meerdere juridische updates). Zie o.a. <https://connectontech.bakermckenzie.com/eu-ai-act-published-dates-for-action/> en <https://www.dlapiper.com/en-us/insights/publications/2025/08/latest-wave-of-obligations-under-the-eu-ai-act-take-effect>

Standaarden & frameworks (beheer, risico, assurance)

1. **ISO/IEC 42001:2023** – *Information technology – Artificial intelligence – Management system*. ISO. <https://www.iso.org/standard/42001>
2. **ISO/IEC 23894:2023** – *Artificial intelligence – Guidance on risk management*. ISO. <https://www.iso.org/standard/77304.html>
3. **ISO 31000:2018** – *Risk management – Guidelines*. ISO. <https://www.iso.org/standard/65694.html>
4. **NIST AI RMF 1.0** – *Artificial Intelligence Risk Management Framework*. NIST (26 januari 2023). <https://www.nist.gov/itl/ai-risk-management-framework> (PDF: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>)
5. **COSO ERM (2017)** – *Enterprise Risk Management—Integrating with Strategy & Performance*. COSO. <https://www.coso.org/enterprise-risk-management>
6. **IIA (2024)** – *Global Internal Audit Standards*. The Institute of Internal Auditors. <https://www.theiia.org/en/standards/2024-standards/global-internal-audit-standards/>

Principes & guidance

1. **OECD (2019, geactualiseerd 2024)** – *Recommendation of the Council on Artificial Intelligence (AI Principles)*. OECD. <https://oecd.ai/assets/files/OECD-LEGAL-0449-en.pdf> (Overzicht: <https://oecd.ai/en/ai-principles>)
2. **EU High-Level Expert Group on AI (2019)** – *Ethics Guidelines for Trustworthy AI*. Europese Commissie. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>