

1b Vragenlijst Bedrijfskundig & Informatiekundig (organisatie + informatieketen)

1. Strategische uitlijning & portfolio

1. Onze AI-roadmap is aantoonbaar gekoppeld aan organisatiedoelen via een gevalideerde benefits-map.
2. Voor AI-initiatieven zijn expliciete stop/door-criteria en exit-paden vastgelegd in het portfolioproses.
3. Hypotheses achter AI-use-cases worden vooraf geformuleerd en periodiek getoetst (A/B of quasi-experimenteel).
4. We hanteren een value-kanban of vergelijkbaar ritme om strategische prioriteit te bewaken.
5. Beslissingen over opschaling volgen vooraf gedefinieerde drempelwaarden (waarde, risico, impact).
6. De verhouding tussen kernidentiteit en vernieuwing is expliciet gemaakt in besluitnotities.

2. Data-governance & datakwaliteit

1. Kritieke datasets hebben een eigenaar, steward, kwaliteitseisen en lineage-overzicht.
2. Dataminimalisatie, rechtmatige grondslag en bewaartermijnen zijn aantoonbaar geborgd per use-case.
3. We meten en remediëren datakwaliteit (compleetheid, nauwkeurigheid, actualiteit, consistentie).
4. Waardeketens hebben datadeals en DPIA's waar nodig; resultaten zijn traceerbaar.
5. Trainings-, validatie- en productiedata zijn herleidbaar en versieerbaar.
6. Toegang tot data volgt least-privilege met periodieke recertificering.

3. Architectuur & interoperabiliteit

1. AI-componenten volgen vastgelegde referentiearchitectuur en integreren via gestandaardiseerde API-contracten.
2. Model- en feature-stores zijn herbruikbaar over domeinen; metadata is gestandaardiseerd.
3. Artefacten zijn portabel (open standaarden) en vendor-lock-in wordt actief beperkt.
4. Non-functionals (latency, throughput, availability) zijn gedefinieerd en getest per ketenstap.
5. Prompt-gateways/LLM-proxies hanteren beleid voor PII-masking en logging.

6. We gebruiken SBOM's voor AI-pakketten en controleren supply-chain-risico's.

4. MLOps & levenscyclusbeheer

1. Elk model heeft een Model Card met doelgrenzen, datascope en verantwoordelijke beslissers.
2. Reproduceerbare pipelines (CI/CD) met versiebeheer voor data, code en modellen zijn operationeel.
3. Monitoring dekt prestatie, fairness, toxiciteit, hallucinaties en drift met geautomatiseerde alerts.
4. Rollbacks/blue-green of canary releases zijn standaard bij modeluitrol.
5. Human-in-the-loop en override-criteria zijn functioneel en getest.
6. Red teaming, prompt-injectie- en adversarial tests zijn periodiek en leiden aantoonbaar tot fixes.

5. Informatiebeveiliging & privacy

1. Model- en prompt-logs worden beveiligd en voldoen aan privacy-eisen (PIA/DPIA waar nodig).
2. Gevoelige inferentie draait binnen afgesproken jurisdicties; datadoorgifte is beoordeeld.
3. Geheimhouding van bedrijfsgevoelige context wordt technisch afgedwongen bij generatieve AI.
4. Geheimenbeheer (keys, tokens) is centraal en voldoet aan rotation-beleid.
5. Continuïteitsplannen (BCP/DR) dekken model-stores en feature-pipelines.
6. Leveranciers voldoen aan due diligence inclusief pentests en compliance-attesten.

6. Waarde, TCO & duurzaamheid

1. Per use-case zijn business-outcomes en TCO vooraf gekwantificeerd (incl. exploitatie).
2. CO₂/energie-impact van training en inferentie wordt geraamd, gevolgd en geoptimaliseerd.
3. We kiezen waar passend voor efficiëntere modellen/architecturen bij gelijkwaardige kwaliteit.
4. KPI's (waarde, risico, ethiek, duurzaamheid) sturen periodiek portfoliobesluiten.
5. "AI-theater" wordt actief ontmoedigd; projecten zonder bewijs van waarde worden tijdig gestopt.
6. Service-recoverykosten bij AI-fouten zijn ingeschat en belegd.

7. Ethiek-operationalisering (ethics-ops)

1. Ethische toets (proportionaliteit, legitimiteit) is een gate in het release-proces.

2. Uitlegbaarheid en beperkingen worden zichtbaar gemaakt in de interface (user-facing).
3. Toegankelijkheid (inclusie, laaggeletterdheid) wordt getest voor AI-interacties.
4. Bias-metingen hebben drempelwaarden en remediatie-paden.
5. Klant/burger heeft recht op uitleg, bezwaar en herstel; processen zijn ingericht.
6. Significante toepassingen worden geregistreerd in een openbaar algoritmeregister.

8. Sturing & beslissingsbevoegdheden

1. Decision rights (wie mag ingrijpen, wie accordeert) zijn expliciet en bekend.
2. Audit-feedback en incidentreviews leiden tot aantoonbare proceswijzigingen.
3. Shadow-AI en tool-sprawl worden actief opgespoord en gereguleerd (whitelist/blacklist).
4. Formele en informele sturing versterken elkaar; governance is licht waar het kan, strikt waar het moet.
5. Rapportages zijn leer- en verbetergericht, niet alleen compliance-gedreven.
6. Risicohouding (risk appetite) is gedefinieerd per AI-domein.

9. Procesprestatie & effectiviteit

1. AI reduceert verspilling zonder nieuwe bureaucratie of extra “prompt-werk”.
2. Structurele bottlenecks worden aangepakt voordat automatisering plaatsvindt.
3. Efficiëntieverbeteringen schaden vakmanschap of waardigheid van werk niet.
4. Reflectietijd is geborgd naast release-druk.
5. Verbeteringen worden systematisch geborgd in standaarden en SOP's.
6. Throughput en first-time-right verbeteren aantoonbaar in de doelprocessen.

10. Adoptie & capability-opbouw

1. Er is een adoptieplan met sponsors, key users en vrijgemaakte leertijd.
2. AI-literacy-briefings voor leidinggevenden vinden periodiek plaats.
3. Functieprofielen en leerpaden bevatten AI-capabilities met toetsing.
4. Gebruikerstesten valideren begrip, vertrouwen en juiste inzet van AI-output.
5. Notificatie-hygiëne en focustijd zijn afgesproken tegen techno-stress.
6. Medezeggenschap is vroegtijdig betrokken bij taak-herschikking en scholing.

11. Ecosysteem, data-deling & keten

1. Datadeals met partners zijn helder over eigenaarschap, gebruik en risico-deling.
2. We benutten open protocollen en gedeelde leerdata waar mogelijk.
3. IP- en licentie-afspraken dekken trainingsdata en outputjuridica.
4. Externe partners ervaren ons als voorspelbaar en uitlegbaar in AI-governance.
5. Controle op leveranciers-updates en modelversies is geborgd.
6. Ketenbrede risico's en baten zijn transparant verdeeld.

12. Markt & dienstverlening (klant/burger)

1. Inzichten in klant-/burgerbehoeften sturen selectie en ontwerp van AI-interacties.
2. We zijn transparant wanneer AI wordt ingezet in contactkanalen.
3. “Dark patterns” en manipulatieve optimalisatie worden uitgesloten.
4. Menselijke service vangt AI-tekortkomingen adequaat op.
5. Feedback uit AI-interacties wordt actief verzameld en verwerkt.
6. Klanten/burgers ervaren onze AI als zorgvuldig en uitlegbaar.

René de Baaij, Van systeem naar betekenis.

Hoe AI de onderstroom van professionals, management
leiderschap en cultuur in organisaties verandert.

rene@debaaij.com

ver 111101