

Den AI-native EA-playbook for virksomheder i midtfeltet

*Sådan får du enterprise-arkitektur der lever,
uden at ansætte et enterprise-team.*

Fem trin fra regneark til levende model: konsolider i én typet graf, lad AI-agenter vedligeholde den, kobl den til virkeligheden via MCP, sæt mennesker i review-rollen, og lad governance opstå som biprodukt. Med en konkret plan: tre dage til første indsigt.

Virksomheder mellem 100 og 750 ansatte har fået enterprise-problemerne: hundredvis af systemer, AI-agenter i alle afdelinger, compliance-krav fra EU og en bestyrelse der vil se sammenhæng mellem strategi og IT-kroner. Det de ikke har fået, er enterprise-staben til at håndtere det.

Den klassiske løsning, et EA-værktøj plus dedikerede arkitekter plus årlige kortlægningsprojekter, er bygget til organisationer med sekscifrede værktøjsbudgetter. Og selv dér dør modellen typisk af manuel vedligehold: surveys ingen besvarer, dokumenter ingen opdaterer, et repository ingen stoler på.

Denne playbook beskriver en anden vej i fem trin, og slutter med en plan der kan eksekveres på tre dage.

PRINCIPPET BAG DET HELE

Modellen skal arbejde for organisationen, ikke omvendt.

Du er integrationen

Lav en lille test i morgen tidlig. Spørg din organisation: hvilke systemer har vi, hvad koster de, hvem ejer dem, og hvilke understøtter vores vigtigste kundeproces?

I de fleste virksomheder i midtfeltet findes svaret ikke i et system. Det findes i en person. Typisk én bestemt person, som samler svaret ved at åbne økonomisystemet, SSO-oversigten, tre regneark, et gammelt Visio-diagram og sin egen hukommelse. Personen er integrationen mellem virksomhedens værktøjer. Hver gang ledelsen har et spørgsmål, betaler den personen for at være menneskelig middleware.

Det er ikke en holdbar arkitektur, og det bliver værre af to grunde:

Systemerne bliver flere.

Softwareporteføljen i midtfeltet vokser år for år, hver afdeling køber selv ind, og en betydelig del af licenserne bruges reelt aldrig. Overblikket forældes hurtigere end det bygges.

Agenterne er flyttet ind.

Gartner forudser at 40 procent af enterprise-applikationer har task-specifikke AI-agenter inden udgangen af 2026, op fra under 5 procent i 2025. Agenter læser, skriver og handler i dine systemer. Hvis intet samlet sted ved hvilke agenter der findes og hvad de må, har du fået en ny slags skygge-IT, nu med hænder.

Problemet er ikke at dine folk er for langsomme til at vedligeholde overblikket.

Problemet er at nogen overhovedet skal gøre det manuelt.

Loopet, der får det hele til at hænge sammen

AI-native arkitektur er ikke “et EA-værktøj med en chatbot”. Det er tre bevægelser der forstærker hinanden:

1. Konsolidér: fra dokumenter til én typet model.

Strategi, kapabiliteter, systemer, integrationer, data, processer, risici og AI-use-cases samles som typede entiteter med relationer, ikke som slides og regneark. Typerne er pointen: når “system understøtter kapabilitet” og “agent har adgang til data” er relationer i en graf, kan man stille spørgsmål på tværs. Hvad rammes hvis vi udfaser dette system? Hvilke kapabiliteter har ingen digital understøttelse? Hvor overlapper tre værktøjer på samme opgave?

2. Connect: modellen taler med virkeligheden, begge veje.

Data flyder ind fra kilderne hvor beslutningerne faktisk træffes: SSO- og usage-signaler, repositories, kontrakter, dokumenter. Og modellen serverer selv kontekst ud via Model Context Protocol, standarden der har rundet 97 millioner installationer og allerede kører i produktion hos 78 procent af enterprise AI-teams. Dine egne AI-assistenten spørger modellen om organisationens sandhed, i stedet for at gætte.

3. Agenter: arbejdskraften, der aldrig glemmer at opdatere.

AI-agenter skriver udkast til alt det ingen gider: berige systemkortet, opdage dubletter, markere forældede oplysninger, foreslå klassifikationer, udkaste review-agendaer. Mennesker godkender. Det er loopets motor, og grunden til at modellen stadig er sand om et år.

RENTERS RENTE FOR ORGANISATIONSVIDEN

En bedre model giver agenterne bedre kontekst, agenterne holder modellen frisk, og en frisk model gør hver ny connector og hvert nyt spørgsmål mere værd. I Atlas er loopet ikke et diagram i et whitepaper, det er produktets arkitektur: typed graf i bunden, connectors og MCP i begge retninger, og agenter med Verify-review som standard.

03 · KAPITEL 3

Start scrappy. Fem felter er nok.

Klassisk EA fejler ofte på ambitionen: seks måneders metamodel-design før første indsigt. Playbook-svaret er det modsatte, og det fortjener at blive sagt uden omsvøb: start sjusket, men start i noget der kan vokse.

Fem felter pr. system er nok til at komme i gang: navn, ejer, hvad den bruges til, kritikalitet, og hvilken kapabilitet den understøtter. Det er en overkommelig import fra det regneark du allerede har, og det er nok til at agenterne kan begynde at arbejde: finde dubletter, foreslå relationer, markere huller.

Resten af felterne kommer ikke fra et projekt. De kommer fra driften: hver gang en agent beriger en post og et menneske bekræfter, bliver modellen en anelse dybere. Efter tre måneder har du den metamodel du ellers ville have designet i et halvt år, bare med den forskel at den afspejler virkeligheden.

Modellér aldrig noget på forhånd, som en agent kan udlede og et menneske kan bekræfte senere.

04 · KAPITEL 4

Agentens anatomi (og hvorfor review ikke er valgfrit)

En agent i produktion består af tre ting:

- **Instructions:** hvad den skal, og hvad den ikke må. Formuleret så et menneske kan læse og rette det.
- **Connections:** hvilke kilder og moduler den har adgang til. Ikke "alt", men en eksPLICIT liste.
- **Triggers:** hvornår den kører. På skema, ved hændelse, eller på anmodning.

I Atlas er det anatomen for Architecture Agents: navngivne AI-teammates som Portfolio Librarian (beriger systemkortet), Freshness Warden (finder det forældede), Duplicate Hunter (finder overlap), Compliance Scout (overvåger AI Act-signaler) og Review Clerk (udkaster mødeagendaer). Du kan se hvad hver agent har lavet, hvornår, og på hvilket grundlag.

Og så det vigtigste designvalg i hele playbooken: **agenter skriver aldrig direkte i modellen**. De skriver forslag. Forslagene lander i Verify-indbakken, hvor et menneske bekræfter, retter eller afviser. Det lyder som et beskedent UI-valg. Det er det ikke. Det er forskellen på et system du kan stille til ansvar, og et system du må håbe på.

TOMMELFINGERREGLER, UANSET VÆRKTØJ

En agent, ét ansvar. Alle skrivninger gennem review, indtil en agent har optjent tillid på et smalt, målbart område. Hver agent har en ejer med navn og et formål på skrift; en agent uden ejer er skygge-IT. Og log alt: loggen ER din compliance-evidens den dag nogen spørger.

05 · KAPITEL 5

Kobl til dér, hvor beslutningerne træffes

Rækkefølgen på integrationer afgør om modellen bliver brugt eller beundret. Kobl ikke til efter hvad der er teknisk nemt. Kobl til efter beslutningsvolumen:

- **Identitet og adgang først** (SSO-oversigten): giver det reelle systeminventar og de første usage-signaler, inklusive kandidater til skygge-AI.
- **Økonomisignalerne** (kontrakter, fornyelsesdatoer, licensantal): gør modellen interessant for CFO'en fra uge ét. "Tre værktøjer på samme kapabilitet fornyes inden for 60 dage" er en sætning der betaler for hele øvelsen.
- **Arbejdssystemerne** (repositories, ticketing, dokumentation): føder agenterne med den kontekst der gør deres udkast præcise.
- **AI-værktøjerne selv via MCP:** når jeres assistenter og agenter henter organisationskontekst fra modellen, bliver modellen infrastruktur i stedet for rapportering.

Bemærk skiftet i økonomien: nu hvor SaaS-leverandørerne selv begynder at shippe MCP-servere (Forrester venter 30 procent i 2026), skal man ikke længere bygge N integrationer. Man skal have én klient og et sted at samle konteksten. Det flytter spørgsmålet fra "hvad kan vi nå at integrere?" til "hvor skal sandheden bo?".

06 · KAPITEL 6

Governance fra dag ét, uden at det føles sådan

Midtfeltet har fået en gave, og den er underligt ukendt: Digital Omnibus-pakken, godkendt af Europa-Parlamentet 16. juni 2026, udskød AI Act's high-risk-forpligtelser til 2. december 2027 og udvidede lempelsesregimet til virksomheder med op til 750 ansatte eller 150 millioner euro i omsætning. Forenklet dokumentation, reducerede bøder, adgang til regulatoriske sandkasser.

Oversat: I har 16 måneders runway og et regime designet til jeres størrelse. Panik er ikke længere en strategi, og det er udsættelse heller ikke. Den fornuftige plan er kedelig og overkommelig:

- **Inventar:** registrér jeres AI-systemer og agenter ét sted, med formål og ejer. Det er søjle 1 i al AgentOps-tænkning: man kan ikke styre det, man ikke kan se.
- **Klassifikation:** afgør pr. use-case om den er omfattet, og i hvilken kategori. Dokumentér rationalet, ikke bare konklusionen.
- **Kadence:** sæt review-rytme på porteføljen, så klassifikationer og ejerskab ikke rådner.

I Atlas er det AI-registret og AI Act-cockpittet: porteføljen, forpligtelserne for jeres størrelse, dokumentations-skabelonerne og tidslinjen samlet, med provenance på hvorfor hver vurdering ser ud som den gør. Compliance bliver et biprodukt af at drive porteføljen ordentligt, ikke et årligt projekt med eksterne timer.

07 · KAPITEL 7

Tre dage til første indsigt

Planer der kræver et kvartal, bliver ikke til noget i midtfeltet. Her er den plan vi anbefaler, og som Atlas' onboarding er bygget om:

Dag 1: Importér rodet.

Tag systemlisten som den er: regnearket, SSO-eksporten, økonomilisten. Fem felter pr. system, huller er tilladt. Lad agenterne køre første berigelse og dublet-scan. Resultatet ved fyraften: et systemkort der er 80 procent sandt, hvilket er cirka 80 procentpoint mere end i går.

Dag 2: Kobl de to første kilder.

SSO og økonomisignaler. Verify-indbakken begynder at fyldes med forslag: manglende ejere, sandsynlige dubletter, systemer uden kapabilitet. En time med kaffe og bekræftelses-klik. Ved fyraften: modellen har relationer, ikke bare rækker.

Dag 3: Første agent i arbejde, første view til ledelsen.

Aktivér freshness- og dublet-agenterne på kadence, og åbn boardroom views: Capability Map med system-understøttelse, Value Streams med support-gaps, Portfolio med kritikalitet. Tag det med til næste ledermøde og stil ét spørgsmål: "hvilket af de her huller gør mest ondt?" Så er arkitekturarbejdet i gang, og det startede med værdi i stedet for med metode.

SLUTBILLEDET, NOGLE MÅNEDER SENERE

Mandag morgen ligger ugens brief klar: hvad ændrede sig i landskabet, hvad venter på beslutning, hvilke fornyelser nærmer sig, hvad fandt agenterne af overlap, og hvor er modellen blevet usikker. Arkitekten ordner på en formiddag det, der før var et kvartalsprojekt. Og da bestyrelsen spurgte til AI-porteføljen, tog det ti minutter at trække rapporten, ikke tre uger.

NÆSTE SKRIDT

Se playbooken køre, ikke bare beskrevet.

Atlas har et komplet demo-workspace med en realistisk virksomhed i midtfeltet, inklusive rodet: dubletter, forældede ejere, en skygge-agent og et opkøbs-scenarie. Vi lover ikke guld. Vi lover en model der stadig er sand til foråret.

[Prøv Atlas](#)

[Tal med os om EA](#)

[ATLAS.SPEKIR.COM](#) · [HELLO@SPEKIR.COM](#)

OM SPEKIR

Spekir er et dansk produkthus med ét fokus: at give virksomheder i midtfeltet den arkitektur-kapacitet de store køber sig til med dedikerede teams. Atlas er Spekirs platform: det governede kontekstlag for den agentiske virksomhed.

KILDER

Gartner (26/8-2025): 40% af enterprise-apps får task-specifikke agenter ultimo 2026, fra <5% i 2025. Europa-Parlamentet, Digital Omnibus (16/6-2026): Annex III udskudt til 2/12-2027; lempelsesregime udvidet til 750 ansatte / EUR 150M (Travers Smith, Latham & Watkins, Morgan Lewis, Stibbe, DLA Piper). MCP-adoption 2026: 97M installationer; 78% af enterprise AI-teams i produktion; Forrester: 30% af SaaS-leverandører shipper egne MCP-servere. Gartner CIO Agenda 2026: 33% beviser konsekvent finansielle AI-outcomes; AI-spend ca. +35% YoY.