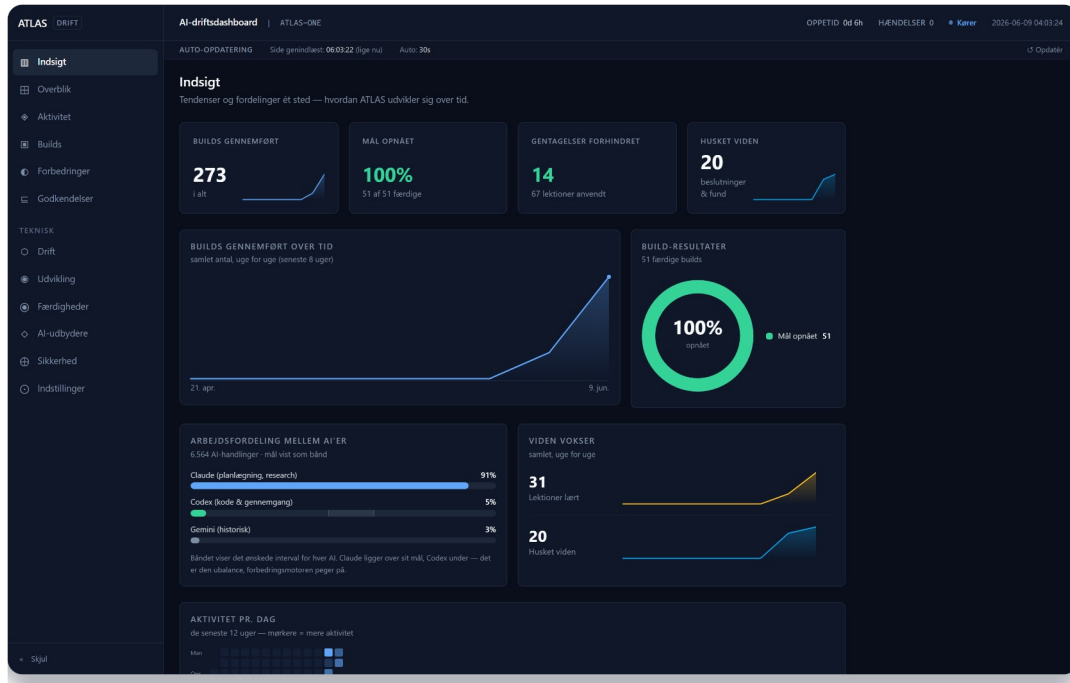


PRODUKT-SHOWCASE

# ATLAS-ONE

*Et selvudviklende AI-operativsystem til autonom softwareudvikling*



Live driftsdashboard: [atlas-runtime.knowision.com](https://atlas-runtime.knowision.com) · Claude + Codex (+ Gemini failover) · 13.750+ signerede events · fem produkter i drift

## Hvad ATLAS-ONE er

---

ATLAS-ONE er ikke en AI-assistent, man stiller spørgsmål. Det er et operativsystem til autonom softwareudvikling — et runtime, hvor AI-modeller ikke bare skriver kode, men driver hele udviklingsforløbet selv: planlægger, researcher, bygger, reviewer sit eget arbejde adversarielt, validerer mod virkeligheden og lukker først en opgave, når målet beviseligt er nået. Mennesket sætter retningen; systemet ejer eksekveringen — på tværs af mange skridt, uden at skulle spørge om lov ved hvert eneste.

Det centrale, og det der adskiller ATLAS-ONE fra et almindeligt AI-værktøj, er at systemet udvikler sig selv. Det overvåger løbende AI-landskabet på internettet — nye modeller, nye frameworks, nye teknikker — og holder fundene op mod sine egne komponenter: Hvad findes der nu, som ville gøre min hukommelse, min routing eller min verifikation bedre? Når det finder noget brugbart, omsætter det fundet til en konkret forbedring, implementerer den bag sine egne sikkerheds- og kvalitetsgates, og lægger den i drift. Systemet bliver målbart bedre for hver fejl det laver og hver opdagelse det gør — ikke fordi en udvikler bygger videre på det, men fordi det bygger videre på sig selv.

Den samme motor peger udad mod produkterne. ATLAS-ONE driver i dag fem selvstændige produkter (bl.a. et AI/world-overvågningssystem, et GDPR/AI Act-compliance-værktøj og en marketing-research-motor), og den samme research-, vurderings- og implementeringsløkke, der forbedrer platformen, kan researche markedet og forbedre produkterne. Det er på den måde ét system, der både bygger sig selv og bygger forretning oven på sig selv.

Hele systemet er styret af en mission-doktrin: en opgave lykkes på sit mål, ikke på at testene er grønne. “Tests passerer” accepteres eksplicit som nødvendigt, men aldrig tilstrækkeligt — der kræves bevis på at tingen faktisk virker i drift. Og alt er indespærret i én sikkerhedsgrænse: systemet kan køre dusinvis af skridt autonomt, men stopper og spørger ved den genuint farlige delmængde (irreversibel infrastruktur, hemmeligheder, destruktive sletninger, udgivelse).

## Hovedkomponenter

---

Hver komponent er beskrevet ved hvad den gør — med et konkret eksempel, ikke et løfte.

### Memory — strategisk hukommelse

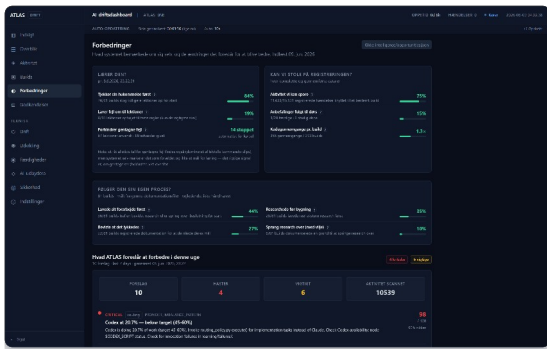
Beslutninger, mønstre og lessons gemmes på tværs af sessioner og genkaldes relevans-rangeret.

**Eksempel:** *før en opgave løses, henter systemet de mest relevante tidligere erfaringer frem — så samme fejl ikke begås to gange, selv måneder senere.*

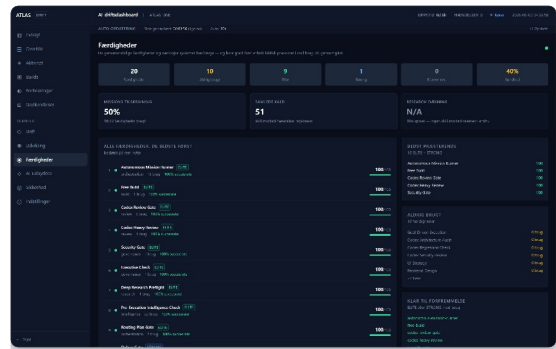
### Learning — selvforbedrende loop

Hver uventet fejl skrives ned, rettelsen bliver en “lesson”, og beviste, gentagne mønstre forfremmes automatisk til en pre-action guard.

**Eksempel:** *efter at have ramt samme cross-shell-fejl nok gange, blokerer systemet nu kommandoen før den køres — fejlen kan ikke længere ske.*



LÆRINGS-METRIKKER: HUKOMMELSE, LEKTIONER, FORHINDREDE GENTAGELSER

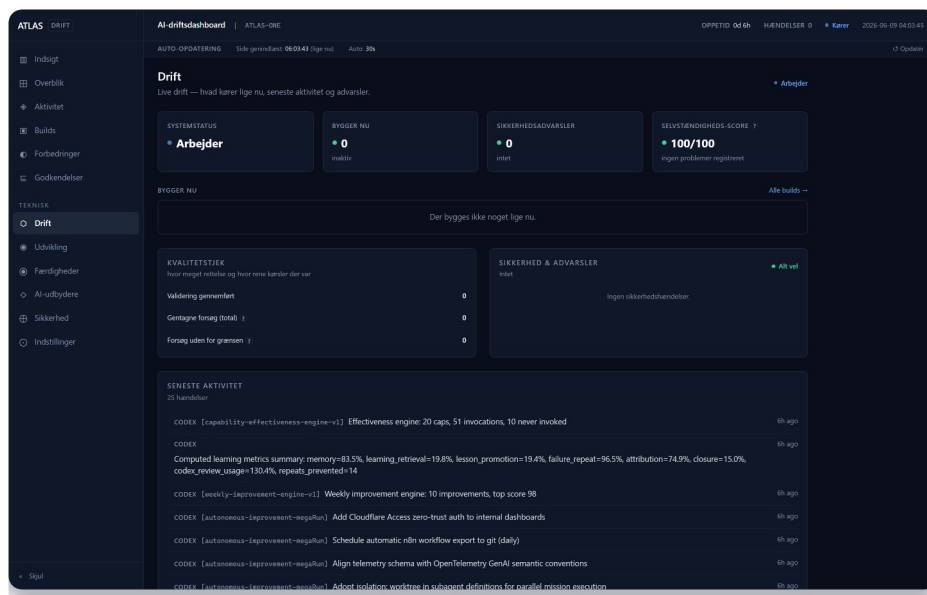


FÆRDIGHEDER RANGERET EFTER EFFEKTIVITET

## Telemetry — auditbar single-writer-log

Hver væsentlig handling logges som et signeret event (13.750+ til dato).

**Eksempel:** *bevis-events er HMAC-signerede, så en mission ikke kan “forfalske” at et review eller en validering fandt sted.*

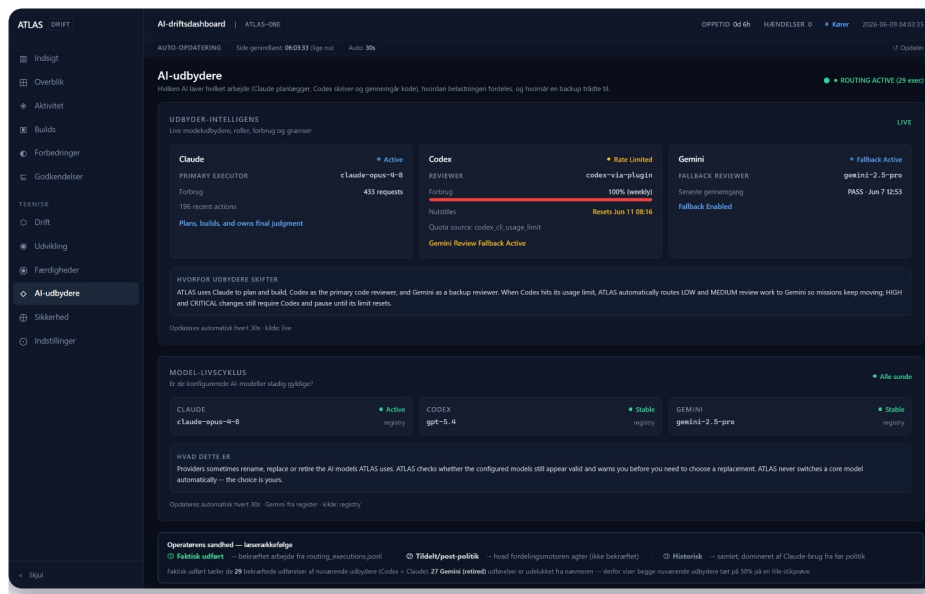


DRIFT: KVALITETSTJEK, SIKKERHED OG HÆNDELSESLOG

## Routing — risiko-tieret multi-model

Opgaver dirigeres til rette model og omkostningsniveau efter risiko, med automatisk failover.

**Eksempel:** *en høj-risiko sikkerhedsændring tvinges igennem dyrt adversarisk Codex-review; er Codex udtømt, overtager Gemini provisorisk så natkørslen ikke stopper.*

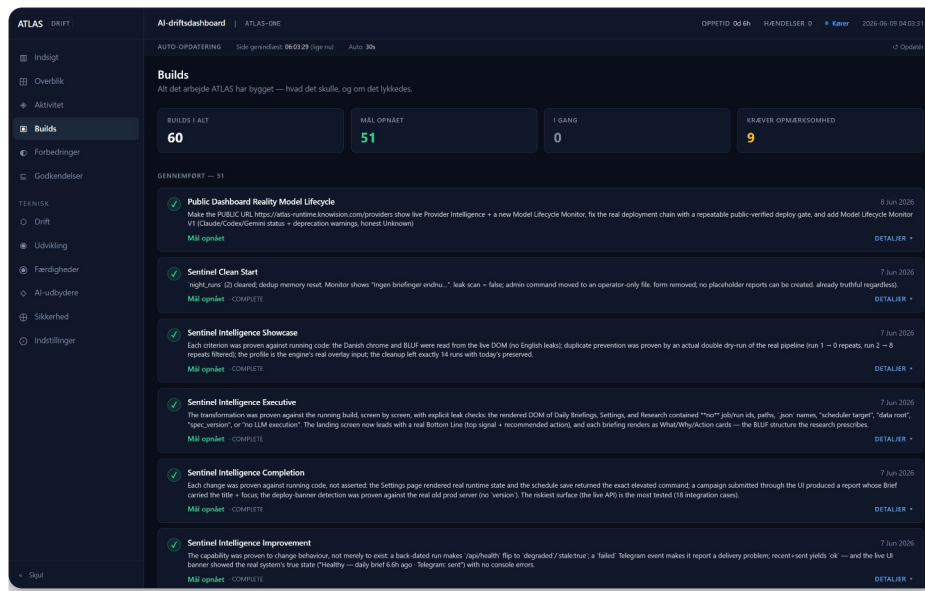


AI-UDBYDERE, ROUTING-BESLUTNINGER OG MODEL-LIVSCYKLUS

## Closure / Verification — fire uafhængige “oracles”

Goal-closure, grounding, trajectory og Independent Goal Verification afviser falsk “færdig”.

**Eksempel:** IGV-oraklet afviser “vi implementerede X / filen findes / testen passerer” som bevis og kræver dokumentation af hvad systemet nu faktisk gør.



BUILDS MED RESULTAT: MÅL OPNÅET, DELVIST ELLER IKKE AFSLUTTET

## Research — preflight + faithfulness

Større opgaver kræver internet-research før implementering, og påstande tjekkes mod deres kilder.

**Eksempel:** faithfulness-oraklet flagger et research-fund, der ikke kan spores til en citeret kilde.

## Continuity — durable handoff

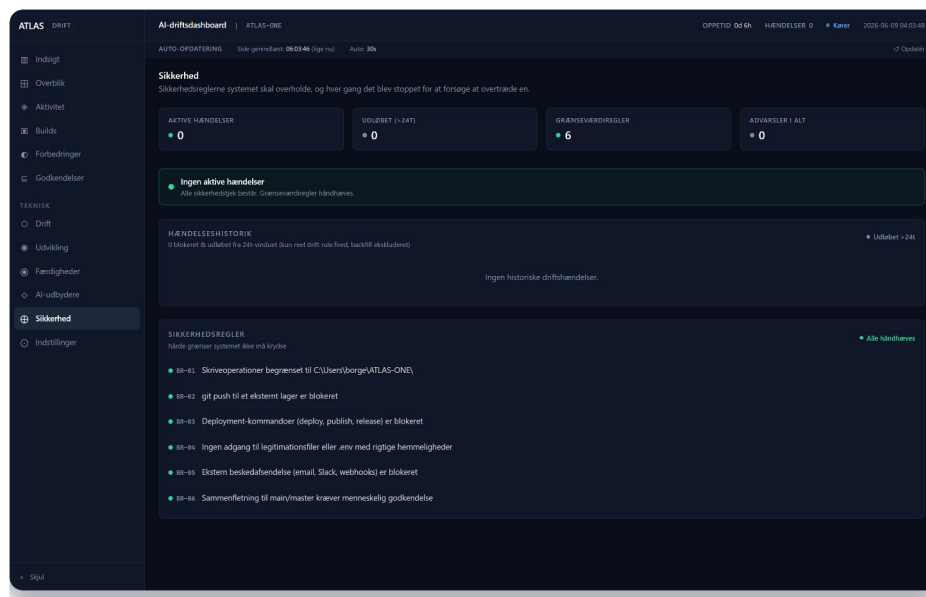
Hver tilstandsændring gemmes i et SQLite-checkpoint, der valideres som genoptageligt.

**Eksempel:** *hvis sessionen crasher eller context-vinduet fyldes, kan en frisk session genoptage præcis hvor den slap — ikke starte forfra.*

## Security — boundary + supply-chain

Alt er scopet til én disk-boundary, og eksterne repos neutraliseres ved kloning.

**Eksempel:** *et klonet repos skills/plugins/hooks sættes i karantæne i tre forsvarslag, så det ikke kan auto-wire sig ind i en live session.*

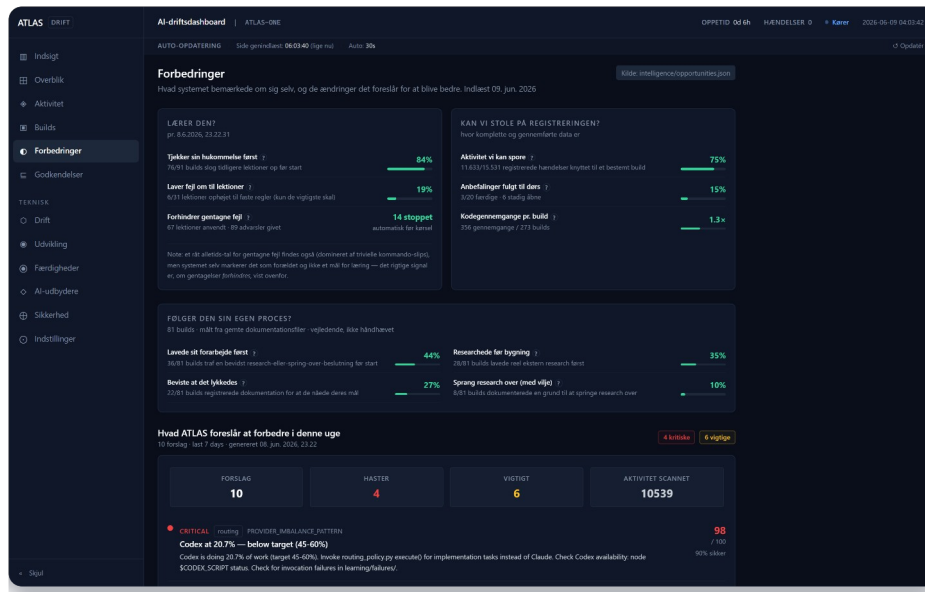


SIKKERHEDSGRÆNSE-REGLER OG HÆNDELSHISTORIK

## Intelligence — ugentlige beslutningsmotorer

Motorer der læser systemets egne data og foreslår næste skridt: kampagne-anbefaler, executive-brief, failure-risk-brief, opportunity-finder.

**Eksempel:** *en ugentlig motor udpeger den højest-værdi sikre opgave at tage fat på næste gang — fra systemets egen telemetry.*



FORBEDRINGER: UGENTLIG MOTOR OG OPPORTUNITY-RADAR

## Guards — runtime-sikkerhedsnet

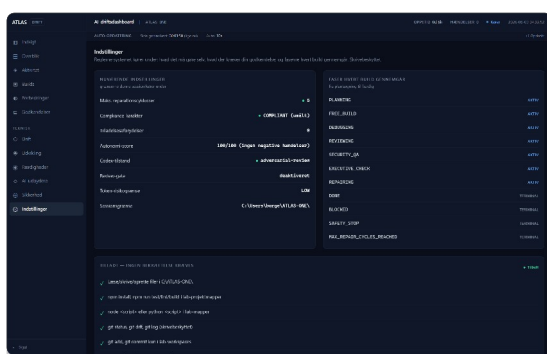
Stuck-loop-detektion og failure-prevention kører på hvert tool-kald.

**Eksempel:** *gentager systemet samme fejlende handling 5 gange, eskaleres en advarsel og missionen markeres BLOCKED frem for at loope i det uendelige.*

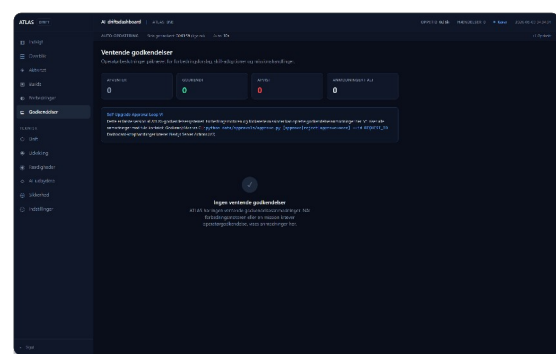
## Governance — doktrin + 26 regler + 12 skills

Et lag der definerer hvad systemet må gøre uden menneskelig godkendelse, og hvornår det skal stoppe.

**Eksempel:** *inden for mandat kører det selv gennem dusinvis af sikre skridt; ved git push, hemmeligheder eller destruktive sletninger stopper det og spørger.*



AUTONOMI-REGLER OG MISSIONS-TILSTANDE

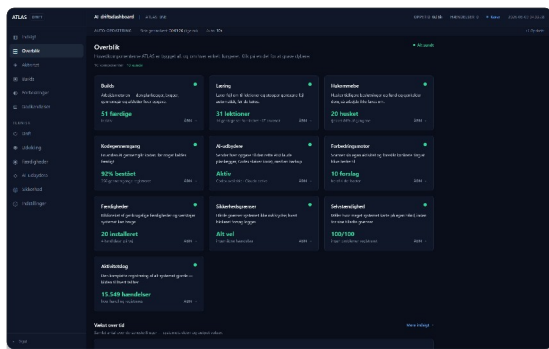


GODKENDELSE: DET FARLIGE SÆTTES I KØ TIL Mennesket

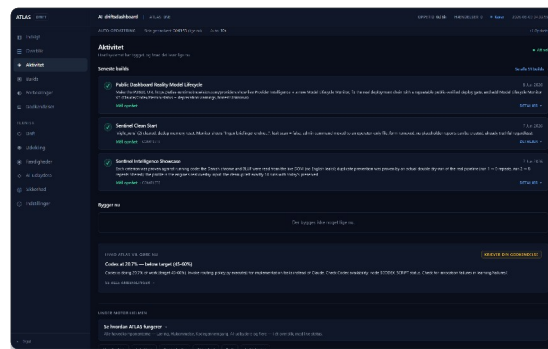
## Dashboard — live offentligt provenue

Et offentligt dashboard viser providers, model-livscyklus og kørselsstatus live.

**Eksempel:** *en simuleret model-deprecation udløste en synlig offentlig alarm på atlas-runtime.knowision.com.*



OVERBLIK: KOMPONENT-SUNDHED PÅ TVÆRS AF SYSTEMET



AKTIVITET: SENESTE BUILDS OG NÆSTE SKRIDT

## Hvor jeg ser det gå hen

Retningen følger den retning AI selv bevæger sig i: fra assistenter, der svarer, til agenter, der arbejder — og det næste skridt er agenter, man kan stole på uden opsyn. Det er præcis det problem ATLAS-ONE er bygget omkring. Tre udviklingslinjer ligger lige for:

- **Fra én autonom bygger til en flåde.** Arkitekturen kører i dag mest én orkestreret tråd. Det naturlige næste skridt er flere parallelle, specialiserede agenter, der bygger forskellige dele samtidig og konvergerer mod ét review — governance-laget til det findes allerede.
- **Selvudvikling, der lukker hele cirklen.** I dag foreslår og implementerer systemet forbedringer bag menneskelige gates. Målet er en stramt afgrænset, fuldt lukket løkke: opdag → vurdér → byg → bevis → idriftsæt → mål effekten → lær — hvor mennesket flytter sig op til kun at sætte retningen og godkende det farlige.
- **Fra værktøj til platform.** Det samme runtime, der bygger sig selv, kan hoste og forbedre vilkårligt mange produkter. Med modeller der bliver billigere og dygtigere kvartal for kvartal, bliver “en autonom udviklingsafdeling i en mappe” en realistisk retning — ikke en metafor.

## Bygget under begrænsning — og hvad der venter uden

ATLAS-ONE er bygget på almindelig hardware og et stramt budget. Den begrænsning har formet arkitekturen til det bedre: fordi jeg ikke kunne køre tunge lokale modeller eller betale for ubegrænset compute, blev systemet tvunget til at være omkostningsbevidst og opfindsomt frem for råstærkt. Routing efter risiko (dyrt review kun hvor det betaler sig), provider-failover når en kvote rammer loftet, durable checkpoints så ingen kørsel spildes, og en doktrin der måler bevist værdi frem for forbrugt tid — det er alt sammen designsvar på “gør mere med mindre”. Begrænsningen er, med andre ord, blevet til en disciplin, der ville gavne ethvert produktionssystem.

Var jeg ikke begrænset af hardware og økonomi, ville den næste fase være ligetil: køre den autonome flåde 24/7 i skala i stedet for i afmålte natkørsler; supplere med lokale modeller til de billige, høj-volumen-opgaver, så de dyre frontier-modeller spares til det svære; og lade selvudviklings-løkken køre kontinuerligt frem for i bursts. Fundamentet er allerede bygget til det — det der mangler, er ikke arkitektur, men compute. Det jeg har vist her, er at designet holder; det der står tilbage, er at skrue op for kraften.