

Hur skiljer sig ett AI projekt från ett traditionellt projekt?

AI projekt skiljer sig ofta åt från andra projekt, som t.ex. IT app projekt. App projekt är lösningsspecifik. När det blir svårt att specificera en lösning blir resultatet osäkert och riskabelt. Sådana projekt faller då under typen top-down programmering.

I motsats följer AI projekt, med Proof of Value (POV), en bottom-up approach. I dessa fall drar AI slutsatser från sina egna regler och processer genom att arbeta med stora dataset.

Utvecklingslandskapet för AI tenderar att öppna upp för flera möjligheter allt eftersom tekniken mognar. Det innebär att projektet, för att det ska bli klassat som färdigt, måste passera flera stadier av utforskande, framgång och försök. Även om utgången av ett sånt här projekt allt som oftast leder till en hög nytta för slutkund, så leder det ändå ofta till höga utvecklingskostnader och förlängda utvecklingstider.

Dela in AI i två olika kategorier

Den första delen i att planera ett AI projekt startat med att en grupp förstår vilken kategori den tillhör. Första kategorin handlar om projekt som är tämligen lika, som att översätta ett språk till ett annat, eller konvertera bilder till ord. Andra kategorin är mer komplex. Den handlar t.ex. om att upptäcka hjärtslag eller monitorera sömn.

De två kategorierna kräver två olika lösningar, tillämpning av en befintlig AI eller skapandet av en skräddarsydd projektledningslösning för AI.

Befintlig AI lösning

Det finns många tillfällen där det blir allt vanligare att inkludera AI. Det innebär att där finns färdigsnickrade lösningar som man behöver för att integrera AI inuti sin lösning. Vanliga plattformar för detta är Microsoft Azure AI, Google AI Platform, och Amazon Machine Learning etc.

Skräddarsydd AI lösning

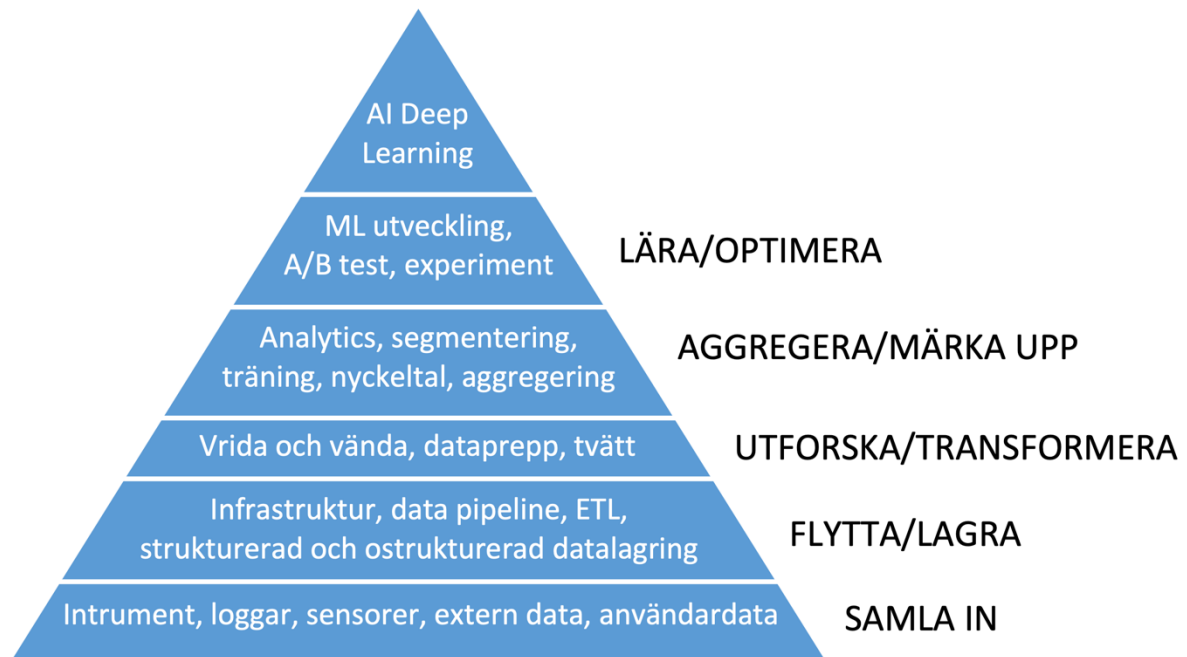
Ifall där är ett mer komplext projekt, exempelvis app driven av neurala nätverk som ger kunder insikter i deras hälsa baserat på deras röst (det finns faktiskt!), måste man titta på skräddarsydd utveckling för AI lösningar.

Nycklar för att nå målet

Många har insett att det viktigaste för att lyckas med sina AI projekt är två saker: människor och data. Bara när man har dessa båda kan AI lyckas med att förbättra användarupplevelsen i grunden. Man börjar med att ta in experter från olika avdelningar som en lösning ägs av eller kan kopplas till, oavsett personernas tekniska kunnande. Det är nödvändigt att mata en algoritm med domänspecifik data för att göra ett AI system effektivt och objektivt.

Nästa del är data. När man inte lagrar data korrekt eller inte lagrar den i sin helhet blir det oanvändbar. Där finns två olika datatyper som verksamheter genererar, strukturerad data (datum, adresser etc.) och ostrukturerad data (fakturor, ljudupptagningar, epost etc.). I arbete med AI bör man överväga båda datatyper.

Det finns en del steg som man behöver ta igenom sin data så att den kan bli användbar för t.ex. deep learning.



Ju fortare data hittar sin plats i denna pyramid, som är baserad på Maslow's behovstrappa, desto fortare kommer ett AI projekt att börja snurra och desto större blir möjligheterna att specialister får arbeta mer på AI modellering och mindre på data prepp.

Varför misslyckas AI projekt?

Felriktade förväntningar

Väldigt ofta får AI projekt inte se dagens ljus på grund av felriktade förväntningar. Grundproblemet till denna utmaning gällande AI och affär/verksamhet uppstår ofta till följd av höga kortsiktiga förväntningar på en teknologi som i sig själv arbetar långsiktigt.

En annan möjlig felriktad förväntan är att AI lösningar ska vara pricksäkra nog för att möta olika användares upplevelse. T.ex. en förväntan inom streamad musik att "nästa låt" som AI föreslår är exakt vad användaren tycker hör till genren är ett problem. Det är därför som många tjänster ofta använder uttryck som "kanske" (eng: "may") när man visar upp en produkt eller tjänst som deras användare skulle kunna vara intresserade av härnäst.

Ineffektiv datahantering

AI tenderar att fatta fel beslut som en följd av felaktigt dataset. Problemet för AI lösningar i olika projekt uppstår när data är inkorrekt eller inkomplett, alltså inte anpassad för att passa AI modellen.

För att ett AI system ska fungera som förväntat, är det nödvändigt att ha förfinad data som kan användas för att lärande och analysera mönster. För att få till data på ett bra sätt för AI är det bra att fokusera på att dela upp strukturerad och ostrukturerad data.



1. Identifiera problem.

I ett AI projekt börjar vi med att identifiera problemet. Vi fokuserar här på två frågor: *"Vad är det vi verkligen vill lösa?"* och *"Vilket är vårt önskade resultat?"*. Viktigt att förstå är att AI i sig självt inte är en lösning utan ett medel eller verktyg för att möta ett visst behov. Vi skulle kunna välja flera olika lösningar som alla hade varit behjälpta av AI, men som faktiskt inte beror på AI.



2. Testa om lösningen passar

Detta steg svarar förmodligen på hur vi bör gå tillväga i vårt AI projekt. Innan vi påbörjar någon utvecklingsprocess för AI så är det först viktigt att testa och bli säkra på att intressenter är villiga att betala för det de bygger upp. Vi behöver testa om lösningen passar genom att använda oss av olika tekniker. Det kan exempelvis vara en Design Sprint (Google) eller LEAN, för att nämna några.

En fördel med AI teknologier är att vi relativt snabbt kan utforma en enklare variant av en lösning med hjälp av vanliga människor eller en MVP. Fördelen här är inte bara en enkel analys av lösningen utan också att inom rimlig tid garantera att produkten faktiskt behöver en AI lösning.



3. Förbereda och hantera data

Om vi kommer dit att vi vet att där finns en kundbas för vår lösning, och vi är övertygade om att vi kan bygga rätt AI, så startar vi hanteringen av maskininlärningsprojektet genom att samla data och ordna vår datahantering.

Vi börjar med att dela upp datan i strukturerad och ostrukturerad. Denna fas kan bli enklare om vi arbetar med enstaka datasets. När det handlar om flera olika datasets kan det bli knepigare, T.ex. om datan lagras i silos och inte kan sammanföras smidigt kan det bli mycket svårare.



4. Välja rätt algoritim

Här nämns inte alla algoritmer som kan tänkas finnas, men vi tittar utifrån supervised learning och unsupervised learning.



Supervised learning kan enklast delas in i två tekniker, ena är klassificering, andra är regressioner. Enkelt förklarat predikterar klassificering en märkning, medan regression predikterar en kvantitet. Vi använder oftast klassificering när vi vill prediktera sannolikheten för ett visst utfall, t.ex. om det kommer regna imorgon. Med en regression kan vi istället kvantifiera scenariot, t.ex. när vi vill veta om en plats kommer bli översvämmad. Det finns så klart mängder av andra algoritmer att välja mellan beroende på projektets natur.

Unsupervised learning, här kan algoritmerna skilja sig mycket åt eftersom datan inte är organiserad eller enligt ett särskilt format. Vi kanske vill använda klustringsalgoritmer för att gruppera objekt eller associationsalgoritmer för att hitta länkar mellan olika objekt.



5. Algoritmträning

När vi valt vår algoritm vill vi gå vidare med att träna modellen genom att stoppa in data i den. Vi betonar vikten av modellens träffsäkerhet. Vi behöver här arbeta med minsta acceptabla tröskelvärden och vara noga med statistiken. Detta för att föra utvecklingen framåt, och alltså på ett sätt som kräver mindre finjustering längre fram.



6. Driftsättning

Det kan vara bra att använda färdiga plattformar när det är möjligt. Många plattformar underlättar arbetet avsevärt med att driftsätta ett AI projekt. Där kan också finnas tillfällen då det inte lämpar sig med driftsättning i t.ex. en molnmiljö eller en färdig tjänst och då måste man se över hur man på bästa sätt kan driftsätta en algoritm utifrån de lokala förutsättningar som ges.