



AALBORG UNIVERSITET

REGULERING AF KUNSTIG INTELLIGENS I EU, USA OG KINA I ET VARIETIES OF CAPITALISM PERSPEKTIV



Titelblad

Specialetitel: Regulering af kunstig intelligens i EU, USA og Kina i et varieties of capitalism perspektiv

ECTS: 30

Vejleder: Mads Peter Klindt

Institut: Aalborg Universitet, Institut for Politik og Samfund

Udarbejdet af: Emil Boddum, 20172906

Antal normalsider: 53,4

Antal tegn: 128.199

Antal ord: 17.718

Indholdsfortegnelse

Indledning	1
Problemfelt.....	1
Regulering.....	2
Afgrænsning.....	4
Problemformulering.....	5
Relevans.....	5
Teori.....	6
Videnskabsteoretisk position	6
Kritik af videnskabsmetode	7
Historisk institutionalisme	8
Valg af teori	10
Varieties of Capitalism (VoC)	10
Liberal Market Economies (LME).....	12
Finansiering.....	13
Uddannelse.....	13
Organisering.....	14
Konkurrence.....	14
Coordinated Market Economies (CME)	15
Finansiering.....	15
Uddannelse.....	16
Organisering.....	16
Konkurrence.....	17
Sammendrag af VoC idealtyper og institutionelle forskelle.....	18
EU's Variety of Capitalism.....	18

Kinas Variety of Capitalism.....	20
Kritik af teori.....	21
Metode	23
Kvalitativ metode.....	23
Komparativt casedesign	23
Dokumentanalyse.....	24
Dataindsamling	25
Forskningskriterier	26
Styrker og svagheder ved metodevalget	28
Empiri	28
EU's AI Act	29
USA's executive order om kunstig intelligens	30
Kinas 24 retningslinjer for kunstig intelligens.....	31
Hypoteser	32
Hypotese 1	33
Hypotese 2	33
Hypotese 3	33
Hypotese 4	34
Undersøgelsesspørgsmål.....	34
1. Hvordan afspejler regulering af kunstig intelligens de respektive aktørers varieties of capitalism?	34
2. Hvordan påvirker de institutionelle kontekster i koordinerede markedsøkonomier (EU), statskapitalisme (Kina) og liberale markedsøkonomier (USA) tilgangen til AI-regulering?	35
3. Hvad er de specifikke reguleringsrammer for AI-udvikling?	35
Analyse	35

EU	36
1. Regulering og relation til VoC.....	37
2. Institutionelle Kontekster og AI-Regulering	37
3. Specifikke Reguleringsrammer for AI-Udvikling	38
USA.....	39
1. Regulering og relation til VoC.....	39
2. Institutionelle kontekster.....	40
3. Specifikke reguleringsrammer	41
Kina.....	42
1. Regulering og relation til VoC.....	42
2. Institutionelle kontekster.....	43
3. Specifikke reguleringsrammer	44
Opsamling på analyse	45
LME-elementer	46
CME-elementer	47
Diskussion.....	48
Hypotese 1	48
Hypotese 2	49
Hypotese 3	51
Hypotese 4	52
Opsamling.....	53
Konklusion.....	54
Bibliografi	56

Abstract

This thesis examines the regulation of artificial intelligence in the United States, China, and the European Union through the lens of varieties of capitalism theory and historical institutionalism. It argues that the distinct approaches to AI governance in these regions are shaped by the institutional frameworks and coordination mechanisms inherent to their respective capitalist models. The thesis employs a qualitative methodology, analyzing key policy documents to investigate how these models influence the values, interests, and actors involved in AI regulation. The findings reveal that the EU's risk-based approach, the US's focus on innovation, and China's state-led approach are deeply rooted in their respective capitalist systems. This study contributes to the ongoing debate on the relationship between society and AI, offering empirical evidence that supports and extends the VoC theory. Additionally, it highlights the importance of historical context and the dynamic nature of AI in shaping regulatory outcomes, providing valuable insights for policymakers and stakeholders navigating the complex landscape of AI governance to ensure better regulation of AI.

Indledning

Kunstig intelligens (AI) er en hastigt voksende teknologi med potentiale til at transformere samfundet grundlæggende. Dette potentiale har medført en intens global konkurrence om at være førende inden for AI-udvikling og -anvendelse. Samtidig har det skabt et presserende behov for regulering for at håndtere de etiske, sociale og økonomiske udfordringer, som AI medfører.

I de seneste årtier har der været en stigende interesse for, hvordan forskellige landes institutionelle rammer påvirker deres økonomiske resultater og konkurrenceevne. Varieties of Capitalism (VoC) teorien har været central i denne diskussion, idet den argumenterer for, at der er flere forskellige veje til økonomisk succes, og at lande kan trives med forskellige kapitalistiske modeller. VoC-teorien giver os derfor et nyttigt redskab til at forstå, hvordan forskellige lande med forskellige politiske og økonomiske systemer tilgår reguleringen af AI.

Specialet vil undersøge reguleringen af kunstig intelligens (AI) i tre forskellige økonomiske kontekster - USA, Kina og EU - ved hjælp af VoC-teorien. Disse cases er valgt, fordi de repræsenterer forskellige kapitalistiske modeller og har en betydelig indflydelse på den globale AI-udvikling. Ved at sammenligne og kontrastere disse tre cases tilgange til AI-regulering kan man opnå en dybere forståelse af, hvordan forskellige kapitalistiske modeller påvirker reguleringen af denne transformative teknologi.

Problemfelt

Den hastige udvikling inden for kunstig intelligens (AI) har indvarslet en ny æra, hvor teknologi og geopolitik er tæt forbundne, som minder om rumkapløbet under den kolde krig. Regeringer verden over anerkender AI's potentiale til at revolutionere industrier, forsvarssystemer og samfundets struktur. I denne globale konkurrence om teknologisk overlegenhed er AI blevet centralt. Fra Silicon Valley til Beijing er kapløbet om at udvikle den mest avancerede AI en topprioritet. Dette speciale vil undersøge AI's udvikling og den geopolitiske kamp om at mestre denne teknologi, som Ruslands præsident Vladimir Putin har beskrevet som nøglen til at "herske over verden" (Associated Press, 2021).

Kunstig intelligens (AI) har en fascinerende historie, der går tilbage til de tidlige dage af datalogien. Alan Turing (2004) skrev i 1948 manifestet for kunstig intelligens "Intelligent Machinery". Heri teoretiserede han, at for at blive betragtet som intelligent, skulle en maskine være i stand til at efterligne menneskelig adfærd så godt, at den ikke kunne skelnes fra et menneske. I 1955, på en konference på Dartmouth University, foreslog forskere formelt betegnelsen "kunstig intelligens" for første gang (Zhang & Lu, 2021). Dette øjeblik markerede begyndelsen på et helt nyt forskningsfelt, der beskæftiger sig med, hvordan maskiner kan efterligne menneskelig intelligens og udføre intellektuelle aktiviteter. De første AI-programmer var simple, men de støbte et fundament til mere komplekse systemer, der ville blive udviklet i de følgende årtier (Zhang & Lu, 2021).

I de seneste år har AI oplevet en periode med hidtil uset vækst på grund af en række faktorer som tilgængeligheden af store datamængder fra internettet og andre kilder har gjort det muligt for forskere at udvikle mere sofistikerede AI-systemer. Indenfor de seneste 5 år er det nye teknologi som grafikkort (GPU'er) fra virksomheden Nvidia, der har gjort det muligt at behandle data hurtigere og mere effektivt. Forskere har udviklet nye algoritmer, der er mere effektive til at lære fra data og løse komplekse problemer. Dette resulterede i udviklingen af den mest kendte store sprogmodel (LLM) udviklet af OpenAI kaldet ChatGPT. LLM'er er en type AI-system, der kan generere tekst, oversætte sprog, skrive forskellige former for kreativt indhold og besvare spørgsmål (Smith, 2023).

Regulering

Kunstig intelligens (AI) er en hastigt voksende teknologi med potentiale til at transformere samfundet grundlæggende. Dette potentiale har medført en intens global konkurrence om at være førende inden for AI-udvikling og -anvendelse. Samtidig har det skabt et presserende behov for regulering for at håndtere de etiske, sociale og økonomiske udfordringer, som AI medfører. Som påpeget af Paul (2023) er regulering et centralt redskab for at afbøde de negative konsekvenser af teknologiudvikling.

I de seneste år er der kommet en række forslag til regulering af AI fra både nationale og internationale aktører. EU har foreslået en AI Act, der sigter mod at skabe en harmoniseret ramme for AI-regulering i hele EU (European Commission, 2021). EU håber, at reguleringen vil sætte en standard for resten af verden gennem den såkaldte "Brussels Effect", hvor EU's lovgivning oftest

smitter af på andre lande, eksempelvis GDPR på det digitale område (Bradford, 2020). USA har præsenteret en række principper og retningslinjer for ansvarlig AI-udvikling, herunder et fokus på gennemsigtighed, retfærdighed og diskrimination (The White House, 2023). Kina har også lanceret en national AI-strategi, der sigter mod at gøre landet til en global leder inden for AI, samtidig med at der lægges vægt på etiske principper og social ansvarligt (DigiChina, 2017). Kinas 24 retningslinjer om kunstig intelligens er Kinas største regulering af kunstig intelligens til specialets dato: 31/5-2024 (DigiChina, 2023).

Disse forskellige tilgange til AI-regulering rejser en række spørgsmål om, hvilke faktorer der driver og former reguleringen, og hvilke konsekvenser forskellige reguleringsmodeller kan have for innovation, konkurrenceevne og samfundsmæssige værdier.

Et centralt spørgsmål er, hvordan AI vil påvirke de forskellige former for kapitalisme, der findes i forskellige verdensdele. Kapitalisme er ikke en ensartet eller statisk model, men snarere et dynamisk og varieret fænomen, der tilpasser sig lokale forhold og institutioner. Forskellige typer af kapitalisme har forskellige styrker og svagheder, som kan gøre dem mere eller mindre modstandsdygtige over for AI's disruptive potentiale.

Varieties of Capitalism (VoC) teorien, som udviklet af Hall og Soskice (2001), giver et nyttigt redskab til at analysere disse forskelle. VoC-teorien skelner mellem to hovedformer for kapitalisme: liberale markedsøkonomier (LME'er) og koordinerede markedsøkonomier (CME'er). LME'er er kendetegnet ved en høj grad af markedsliberalisering, deregulering, privatisering og konkurrence, mens CME'er er kendetegnet ved en høj grad af koordination, regulering, social dialog og samarbejde mellem aktørerne på markedet. Klitgaard & Roederer-Rynning (2009) påpeger dog, at sondringen mellem de to typer af markedsøkonomier bør nuanceres yderligere.

Derudover har forskere argumenteret for, at AI vil skabe nye former for kapitalisme, såsom "AI capitalism" (Verdegem, 2021), "platform capitalism" (Srnicek, 2017) og "surveillance capitalism" (Zuboff, 2019). Disse nye former for kapitalisme er karakteriseret ved en øget brug af algoritmer og data til at styre og optimere økonomiske processer, hvilket kan have betydelige konsekvenser for både LME'er og CME'er. For LME'er kan AI betyde en øget polarisering mellem dem, der ejer og kontrollerer AI, og dem, der bliver erstattet eller underordnet af dem. For CME'er kan AI betyde en udhuling af den sociale koordination, regulering og dialog, som har været grundlaget for deres succes.

Dette speciale vil undersøge, hvordan reguleringen af AI i tre forskellige økonomiske kontekster - USA, Kina og EU - kan forstås i et VoC-perspektiv. Ved at sammenligne og kontrastere disse tre aktørers tilgange til AI-regulering, vil specialet belyse, hvordan forskellige kapitalistiske modeller påvirker reguleringen af denne transformative teknologi, hvilke værdier og interesser der er på spil, og hvordan dette afspejler de respektive økonomiske systemer.

Afgrænsning

Specialet fokuserer på EU, USA og Kina, da disse aktører repræsenterer forskellige kapitalismemodeller som udfoldes i teoriafsnittet, henholdsvis CME, LME og statskapitalisme). De har en betydelig indflydelse på den globale AI-udvikling. For EU er AI Act valgt som det primære dokument, da det er den mest omfattende og ambitiøse AI-regulering på globalt plan. For USA er præsident Bidens dekret valgt, da den repræsenterer den amerikanske regerings officielle holdning til AI-regulering. For Kina er de 24 retningslinjer for kunstig intelligens, da retningslinjerne er den nyeste form for regulering som giver indsigt i Kinas tilgang til AI-regulering.

EU har en kompleks og overordnet lovgivningsramme for AI med AI-forordningen, der kan overskygge de specifikke variationer i regulering på nationalt plan. En analyse af EU's medlemslande ville kræve et separat speciale. Storbritannien har forladt EU efter Brexit og er i færd med at udvikle sin egen lovgivningsramme for AI, men en analyse af Storbritannien ville kræve en dybdegående undersøgelse af den aktuelle situation og dens udvikling. Storbritannien er en liberal markedsøkonomi (LME) ifølge Varieties of Capitalism (VoC), ligesom USA, derfor kunne en sammenligning af disse to lande være interessant, men det er uden for specialets ramme. Denne afgrænsning er foretaget for at sikre en dybdegående og fokuseret analyse af problemstillingen. Ved at sammenligne de tre valgte aktører med forskellige varieties of capitalism kan specialet belyse de komplekse samspil mellem institutioner, koordinationsmekanismer og reguleringen af AI. Denne afgrænsning betyder ikke, at andre aktører ikke er relevante for problemstillingen. Specialets konklusioner kan have implikationer for andre lande med lignende VoC-idealtyper. Tidsmæssigt afgrænser specialet sig til vedtagelsen af EU's AI-forordning, der blev vedtaget med enstemmighed den 2. februar 2024.

Problemformulering

Problemformulering fokuserer på at undersøge, hvordan reguleringen af kunstig intelligens i tre forskellige økonomiske kontekster – USA, Kina og EU – kan forstås i et VoC-perspektiv. Disse cases er valgt, fordi de repræsenterer forskellige kapitalistiske modeller og har en betydelig indflydelse på den globale AI-udvikling. Ved at sammenligne og kontrastere disse tre tilgange til AI-regulering kan vi opnå en dybere forståelse af, hvordan forskellige kapitalistiske modeller påvirker reguleringen af denne transformative teknologi.

Specialets problemformulering:

“Hvordan kan regulering af kunstig intelligens i USA, Kina og EU forstås i et varieties of capitalism perspektiv?”

Denne problemformulering giver mulighed for en dybdegående undersøgelse af, hvordan de politiske, økonomiske og kulturelle faktorer, der er forbundet med forskellige kapitalismemodeller, påvirker reguleringen af AI i disse tre magtfulde markedsøkonomier. Specialet vil undersøge, hvordan de forskellige kapitalismemodeller påvirker reguleringen af AI, hvilke værdier og interesser der er på spil, og hvordan dette afspejler de respektive økonomiske systemer. Derudover vil specialet vurdere de potentielle risici og fordele ved forskellige tilgange til regulering af AI og diskutere implikationerne for det internationale samarbejde og konkurrencen inden for AI.

Relevans

Specialets relevans ligger i dets bidrag til at skabe en bedre forståelse af, hvordan forskellige variationer af kapitalisme kan forme udviklingen og reguleringen af AI. Derudover vil det vurdere de potentielle risici og fordele ved forskellige tilgange til regulering af AI. Specialet er relevant, da det adresserer en aktuel problemstilling med betydelige konsekvenser for samfundet, økonomien og miljøet. Ved at undersøge reguleringen af AI i et komparativt perspektiv, kan specialet bidrage til at informere den politiske debat og beslutningstagning om, hvordan vi bedst kan håndtere de udfordringer og muligheder, som AI-teknologien bringer med sig.

Samtidig er der bekymringer om de potentielle negative konsekvenser af AI, såsom tab af job, øget ulighed og sikkerhedsrisici. EU, USA og Kina forsøger alle at udnytte potentialet af AI-teknologien til at øge innovation, produktivitet og velstand. Hver især har en unik variation af kapitalisme, der kan påvirke, hvordan AI udvikles og reguleres. Komparativ litteratur har ikke anvendt VoC-perspektivet på regulering af kunstig intelligens, hvilket specialet vil bidrage med.

Teori

Dette afsnit præsenterer de teoretiske rammer for specialet, der undersøger reguleringen af kunstig intelligens (AI) i EU, USA og Kina. Specialet tager udgangspunkt i en kritisk realistisk tilgang, der anerkender eksistensen af en objektiv virkelighed, samtidig med at den erkender, at vores forståelse af denne virkelighed er begrænset af vores sociale og historiske kontekst. Denne position giver mulighed for at undersøge ikke blot de observerbare fænomener inden for AI-regulering, men også de dybereliggende mekanismer og strukturer, der påvirker den.

Videnskabsteoretisk position

Specialet tager udgangspunkt i kritisk realisme, en retning, der anerkender eksistensen af en objektiv virkelighed, samtidig med at den erkender, at vores forståelse af denne virkelighed er begrænset af vores sociale og historiske kontekst. Denne position giver mulighed for at undersøge ikke blot de observerbare fænomener inden for AI-regulering, men også de dybereliggende mekanismer og strukturer, der påvirker den.

Kritisk realisme er en filosofisk retning, der forsøger at forklare, hvordan verden er uafhængig af vores opfattelse af den. Kritisk realisme hævder, at der findes en objektiv virkelighed, som vi kan erkende gennem videnskab, men at vores erkendelse er begrænset af vores sociale og historiske kontekst. Kritisk realisme er derfor både realistisk og kritisk over for vores viden om verden (Gorski, 2013).

Videnskabsteori har to centrale begreber: ontologi og epistemologi. Ontologi handler om, hvad der eksisterer, og hvordan det eksisterer. Epistemologi handler om, hvordan vi kan få viden om det, der eksisterer. Kritisk realisme har en stratificeret ontologi, som betyder, at verden består af forskellige lag eller niveauer af virkelighed. Det nederste lag er det reale, som er de grundlæggende mekanismer og strukturer, der bestemmer, hvordan verden fungerer. Det midterste lag er det

faktiske, som er de begivenheder og processer, der sker i verden som følge af det reale. Det øverste lag er det empiriske, som er vores observationer og erfaringer af det faktiske. Kritisk realisme har en fallibilistisk epistemologi, som betyder, at vores viden om verden er usikker og fejlbarlig. Man kan aldrig være helt sikker på, at teorier og forklaringer afspejler det reelle niveau af virkelighed. Man kan kun gøre plausible antagelser baseret på vores empiriske data og vores kritiske refleksion. Kritisk realisme anerkender også, at viden er påvirket af menneskets sociale og historiske position, værdier og interesser, og sprog og kultur. Derfor må man være åbne for at udfordre og revidere viden i lyset af nye beviser og perspektiver (Gorski, 2013).

Kritisk realisme i politologi har til formål at undersøge de underliggende årsager og mekanismer, der driver politiske processer og konflikter. Det betyder, at man ikke kun fokuserer på de synlige og observerbare aspekter af politik, men også på de skjulte og komplekse relationer mellem aktører, institutioner, ideer og interesser. Kritisk realisme forsøger at afdække de strukturelle betingelser og de historiske forandringer, der påvirker politiske handlinger og resultater.

Denne lagdeling muliggør en analyse af AI-regulering, der bevæger sig ud over overfladiske beskrivelser og undersøger de dybereliggende årsager til, hvorfor forskellige lande har valgt forskellige reguleringsmodeller. Ydermere anerkender kritisk realisme, at vores viden er fallibilistisk, hvilket betyder, at den er usikker og kan være fejlbarlig. Dette opfordrer til en kritisk refleksion over de anvendte metoder og resultater og understreger vigtigheden af at være åben for alternative forklaringer og fortolkninger.

Kritisk realisme i politologi er en alternativ tilgang til de mere positivistiske og postmoderne tilgange, der dominerer faget. Positivismen forsøger at anvende naturvidenskabelige metoder til at måle og teste politiske fænomener, men overser ofte de subjektive og normative dimensioner af politik. Postmodernisme udfordrer enhver form for objektivitet og sandhed i politik, men ender ofte med at være relativistisk og skeptisk over for enhver form for forklaring eller kritik. Kritisk realisme tilbyder en mellemvej mellem disse to poler, ved at anerkende både den materielle og den ideelle side af politik, samt både muligheden og nødvendigheden af kritisk analyse (Gorski, 2013).

Kritik af videnskabsmetode

Kritisk realisme hævder at være objektiv, men dens fokus på magt og ideologi gør det vanskeligt at opnå en neutral position. Forskerens subjektive holdninger kan ubevidst påvirke analysen, derfor er tydelig beskrivelse og overholdelse af forskningskriterier afgørende for en objektiv

undersøgelse. Desuden lægger kritisk realisme vægt på kontekst og specifikke situationer, hvilket gør det vanskeligt at generalisere konklusionerne til andre kontekster. Kritisk realismes fokus på fortolkning og forståelse gør det også vanskeligt at falsificere dens teorier. Dette kan føre til dogmatisme og manglende videnskabelig stringens (Gorski, 2013).

Historisk institutionalisme

Historisk institutionalisme er en tilgang til at forstå politiske og sociale fænomener, der lægger vægt på de historiske processer, der former institutioner og aktørers adfærd. Institutioner er de formelle og uformelle regler, normer og mønstre, der styrer menneskelig interaktion i forskellige sammenhænge (Bulmer, 2023). Historisk institutionalisme antager, at institutioner ikke kun er resultatet af rationelle valg, men også af historiske kontingenser, magtkampe og kulturelle påvirkninger. Historisk institutionalisme forsøger at forklare, hvordan institutioner opstår, ændrer sig og påvirker politiske og sociale resultater over tid. Et centralt begreb i HI er stiafhængighed, som refererer til, hvordan tidligere beslutninger og begivenheder skaber "stier", der begrænser fremtidige valgmuligheder. Dette kan føre til, at institutioner og politikker forbliver uændrede, selv når de ikke længere er optimale. Et eksempel på dette er det amerikanske topartisystem, der har vist sig bemærkelsesværdigt modstandsdygtigt over for forandring, på trods af ændringer i vælgeradfærd og demografi. Et andet centralt begreb i historisk institutionalisme er kritiske vendepunkter, hvor institutioner kan ændre sig radikalt som følge af store begivenheder eller kriser. Disse ændringer kan skabe nye stier og åbne op for nye muligheder eller begrænsninger. Et eksempel på et kritisk vendepunkt er den franske revolution, der afskaffede det feudale system og indførte et republikansk styre baseret på menneskerettigheder og national suverænitet. Historisk institutionalisme er en nyttig tilgang til at analysere komplekse og dynamiske fænomener, der ikke kan forklares med simple årsags-virkningsrelationer. Historisk institutionalisme anerkender den historiske kontekst, den institutionelle diversitet og den politiske konflikt som væsentlige faktorer for at forstå samfundets udvikling. Historisk institutionalisme kan forklare, hvordan forskellige typer af kapitalisme har forskellige institutionelle rammer og logikker, som former de politiske aktørers interesser og strategier i forhold til kunstig intelligens (Bulmer, 2023).

Schmidt (2007) argumenterer for, at historisk institutionalisme kan bidrage til at forklare institutionel forandring inden for Varieties of Capitalism. Hun påpeger, at den oprindelige VoC-teori har haft svært ved at redegøre for forandring på grund af dens fokus på institutionel

komplementaritet og positive feedback-effekter, hvilket skaber et billede af systemer i ligevægt, der kun ændrer sig i tilfælde af "punkteret ligevægt" eller revolution.

For at imødekomme denne kritik har nyere VoC-forskning inddraget HI-tilgangen, der fokuserer på evolutionære forandringer, der kan være lige så transformative som kritiske øjeblikke. Denne tilgang erstatter ideen om punkteret ligevægt med inkrementel forandring og historisk determinisme med historisk indeterminisme ved at erstatte stiafhængighed med forskellige processer for stifornyelse, revision eller erstatning.

Schmidt (2007) fremhæver, at institutioner i denne tilgang ses som "regimer" eller "systemer af social interaktion under formel normativ kontrol", hvor aktører følger reglerne, fordi de ikke kun håndhæves, men også er legitime. Institutionel forandring kan ske gennem forskellige processer, der skyldes den løbende interaktion mellem "regelmagere" og "regeltagere", hvilket fører til nye fortolkninger af reglerne. Disse processer kan omfatte "forskydning" og "afhopning", "lagdeling" gennem reform, "drift" gennem forsømmelse, "konvertering" gennem genfortolkning eller omdirigering og "udmattelse" gennem sammenbrud.

I modsætning til den klassiske VoC-teori, der fokuserer på tæt komplementære og koordinerede kapitalistiske systemer, beskriver den reviderede VoC-tilgang komponenterne af løst forbundne, historisk udviklende kapitalismevarianter, der ændrer sig i forskellige hastigheder og på forskellige måder gennem forskellige processer. Denne teoretiske åbning gør det muligt empirisk at kortlægge ændringer i VoC'erne - mod hybridformer og konvergens, såvel som tilbagegang.

Historisk institutionalisme er relevant for specialet, da det kan bidrage til at forklare, hvordan historiske og institutionelle faktorer har formet de nuværende reguleringsrammer for AI i forskellige lande. For eksempel kan man undersøge, hvordan tidligere teknologiske revolutioner har påvirket reguleringen af nye teknologier, og hvordan dette kan give indsigt i, hvordan AI reguleres i dag. Derudover kan HI hjælpe med at forklare, hvorfor nogle lande er mere tilbøjelige til at vedtage en proaktiv tilgang til AI-regulering, mens andre er mere reaktive. Dette kan skyldes historiske erfaringer med regulering af teknologi eller kulturelle forskelle i opfattelsen af risiko og innovation. Ved at inddrage historisk institutionalisme som et supplement til VoC-tilgangen kan specialet opnå en dybere forståelse af de historiske og institutionelle faktorer, der har formet AI-reguleringen i de forskellige lande.

Valg af teori

Varieties of Capitalism (VoC) teorien, som oprindeligt formuleret af Hall og Soskice (2001), har vist sig at være et værdifuldt redskab til at forstå, hvordan forskellige kapitalistiske systemer påvirker økonomiske resultater og virksomhedsstrategier. VoC-teorien fokuserer primært på etablerede industrier og økonomiske sektorer, men dens kernebegreber og brede analytiske rammer kan også anvendes til at analysere reguleringen af nye teknologier som kunstig intelligens (AI).

VoC-tilgangen er brugbar til at besvare specialets problemformulering, da den giver et teoretisk grundlag for at forstå, hvordan forskellige kapitalistiske systemer påvirker reguleringen af kunstig intelligens. Ved at anvende VoC-tilgangen kan specialet undersøge, hvordan de institutionelle rammer og koordinationsmekanismer i liberale markedsøkonomier (LME'er) som USA og koordinerede markedsøkonomier (CME'er) som EU, samt statskapitalistiske økonomier som Kina, påvirker deres tilgang til AI-regulering.

VoC-frameworket giver mulighed for at analysere, hvordan forskellige aktører, såsom staten, virksomheder, fagforeninger og NGO'er, interagerer og påvirker udformningen af AI-regulering. Den giver også mulighed for at undersøge, hvordan forskellige værdier og interesser, såsom økonomisk vækst, social retfærdighed, innovation og kontrol, afspejles i de forskellige reguleringsrammer.

Ydermere kan VoC-tilgangen hjælpe med at belyse, hvordan AI-regulering kan påvirke forskellige aspekter af økonomien og samfundet, såsom innovation, konkurrence, beskæftigelse, ulighed og grundlæggende rettigheder. Ved at anvende VoC-tilgangen kan specialet opnå en dybere forståelse af de komplekse sammenhænge mellem kapitalisme, institutioner og AI-regulering.

Varieties of Capitalism (VoC)

Kapitalisme er et komplekst og dynamisk system, der ikke kan reduceres til en enkelt, universel model. I stedet findes der en række forskellige "Varieties of Capitalism" (VoC), der afspejler de unikke institutionelle og historiske kontekster, som har formet hvert lands specifikke tilgang til markedsøkonomien. Denne teoretiske ramme blev introduceret af Peter A. Hall og David Soskice i bogen *"Varieties of Capitalism: The Institutional Foundations of Comparative Advantage"* (2001), og har haft en betydelig indflydelse på komparativ politisk økonomi og tilbyder et

nuanceret blik på den mangfoldige natur af kapitalistiske systemer. Bogen udkommer på baggrund af en strøm i det intellektuelle miljø efter murens fald, hvor stemmer som Francis Fukuyama (1992) dominerer debatten. Fukuyama argumenterer i sit værk for en global konvergens mod det liberale demokrati, da det liberale demokrati anses som uovertruffen både som økonomisk system og som politisk regime. Bogen forsøger at bidrage med nuancer til denne debat omkring konvergens i kontekst af globalisering.

Det argumenteres i bogen, at der indenfor kapitalismen findes forskellige økonomiske systemer, der fungerer på forskellige præmisser og med forskellige output, men som ikke desto mindre er konkurrencedygtige hver for sig i det globale kapitalistiske system. At forskellige typer af markedsøkonomier har forskellige institutionelle fordele, som påvirker deres specialisering og reaktion på globaliseringen. At institutionerne sætter rammen for hvilke typer af aktiviteter virksomhederne har incitament til at engagere sig i. Dette er VoC-teoriens centrale argument, at forskellige politiske økonomier er præget af distinkte institutionelle infrastrukturer, som virksomheders bestræbelser er delvist afhængige af. Hall og Soskice mener, at institutionel struktur betinger virksomheders strategi, men ikke fuldstændig determinerer den. Denne pointe udfordrer den traditionelle forståelse, hvor virksomheders autonomi var central. VoC vender denne rækkefølge om og fremhæver, at strategi i høj grad er afhængig af struktur. VoC bidrager også med et specifikt fokus på virksomheder som den centrale aktør i den politiske økonomi, dog med øje for rollen som staten, arbejdsgiverforeninger og fagforeninger spiller.

For at uddybe forståelsen af forholdet mellem institutionelle rammer og virksomheders innovative kapacitet, introduceres begrebet institutionelle komparative fordele (Hall og Soskice, 2001). Det påpeges, at specifikke institutionelle infrastrukturer kan styrke virksomheders innovative evner, en faktor afgørende for langsigtet overlevelse. Der skelnes mellem radikale innovationer som udvikling af nye produkter eller arbejdsgange og inkrementelle innovationer som løbende forbedringer af eksisterende produkter eller arbejdsgange. Hall & Soskice (2001) fremhæver, at kortsigtede investeringer, lempelige afskedigelsesregler og generelle fagkundskaber fremmer radikale innovationer som softwareudvikling og bioteknologi. Derfor anses virksomheder i LME'er (Liberal Market Economies) for at være bedre rustet til radikale innovationer sammenlignet med virksomheder i CME'er (Coordinated Market Economies). Omvendt stimulerer

langsigtede investeringer, generøse velfærdspolitikker og specifikke fagkundskaber inkrementelle innovationer. Dette giver CME'er komparative fordele inden for produktmarkeder præget af inkrementelle innovationer i bilindustrien, kemiindustrien og til maskinproduktion. Det antyder også, at liberale markedsøkonomier vil søge at deregulere deres markeder for at tilpasse sig globaliseringen, mens koordinerede markedsøkonomier vil forsøge at bevare eller tilpasse deres institutioner for at beskytte deres konkurrenceevne (Hall og Soskice, 2001, s. 37-41).

VoC-tilgangen er aktør-centreret, hvilket vil sige, at politisk økonomi anses som et terræn befolket af flere aktører, som hver især søger at fremme sine interesser på en rationel måde i strategisk interaktion med andre. De relevante aktører kan være individer, virksomheder eller regeringer. Virksomheder betragtes som de afgørende aktører i en kapitalistisk økonomi. De er de nøgleaktører, der tilpasser sig i mødet med teknologisk forandring eller international konkurrence, hvis aktiviteter summerer sig til det samlede økonomiske resultat (Hall og Soskice, 2001, s. 6)

VoC lægger stor vægt på de strategiske interaktioner, der finder sted mellem forskellige aktører i den politiske økonomi (Hall og Soskice, 2001, s. 5). Institutioner spiller en afgørende rolle for disse interaktioners karakter og dermed for de politiske og økonomiske resultater, der opnås. Kvaliteten af interaktionerne afhænger af den specifikke institutionelle infrastruktur, og når virksomheder koordinerer mere effektivt, forbedres både deres egen præstation og den generelle økonomiske performance (Hall og Soskice, 2001, s. 45).

Varieties of capitalism skildrer to hovedtyper af kapitalisme: den liberale markedsøkonomi (LME) og den koordinerende markedsøkonomi (CME). Det er dog vigtigt at bemærke, at disse modeller er idealtyper og ikke nødvendigvis repræsenterer en perfekt beskrivelse af et specifikt land. I praksis findes der hybrider og variationer af disse modeller, der afspejler unikke nationale kontekster.

Liberal Market Economies (LME)

I liberale markedsøkonomier (LME'er) koordinerer virksomheder primært gennem markedsmekanismer og hierarkier. De er kendetegnet ved fleksible arbejdsmarkeder, lav grad af regulering og stærk konkurrence. LME'er, såsom USA og Storbritannien, fremmer radikal innovation og har en tendens til at have mindre involvering, da staten spiller en begrænset rolle i og sjældent griber ind i markedets funktion. Reguleringen er ofte mindre omfattende og fokuserer på at sikre fair konkurrence og beskytte forbrugernes interesser. Markedsrelationer er kendetegnet

ved udveksling af varer eller tjenesteydelser på armlængdes afstand i en kontekst af konkurrence og formel kontrahering. Som reaktion på de prissignaler, som sådanne markeder genererer, tilpasser aktørerne deres villighed til at tilbyde og efterspørge varer eller tjenesteydelser, ofte på grundlag af de marginale beregninger, som neoklassisk økonomi fremhæver. I mange henseender giver markedsinstitutioner et meget effektivt middel til at koordinere de økonomiske aktørers bestræbelser (Hall og Soskice, 2001, s. 8.)

Finansiering

Virksomheder i LME'er er ofte karakteriseret ved en hierarkisk struktur, hvor beslutninger træffes centralt af topledelsen. Aktionærernes interesser er i fokus, og virksomhedernes primære mål er at maksimere profit og aktiekurs. Virksomheder i LME'er er ofte afhængige af kortsigtet finansiering og aktionærernes interesser. Denne kortsigtede orientering kan dog også have negative konsekvenser, såsom en tendens til at prioritere hurtige gevinster frem for langsigtede investeringer i forskning og udvikling (Hall & Soskice, 2001, s. 27-28).

Uddannelse

Uddannelsessystemet i LME'er er markedsorienteret og fokuserer primært på at give individer generelle færdigheder, der kan anvendes på tværs af forskellige brancher og virksomheder. Dette står i kontrast til CME'er, hvor der ofte er et større fokus på erhvervsuddannelse og branchespecifikke færdigheder. Virksomheders følsomhed over for den aktuelle værdi kræver evnen til at justere produktionen. Arbejdstagerne står derfor over for et flydende arbejdsmarked med korte ansættelsesperioder. Dette fokus på karriere motiverer arbejdstagere til at tilegne sig generelle færdigheder gennem uddannelse, som kan anvendes i forskellige virksomheder. Uddannelsessystemet tilbyder derfor i højere grad formel uddannelse med fokus på generelle færdigheder. Arbejdstagere tager typisk en universitetsuddannelse, hvor de opnår en bred almen viden. Denne type uddannelse reducerer omkostningerne til videreuddannelse, hvilket motiverer virksomheder til at gøre brug af intern uddannelse. Denne videreuddannelse fokuserer ofte på omsættelige færdigheder, da arbejdstagere skal have incitament til at tilegne sig dem. Resultatet er en arbejdsstyrke med høje generelle færdigheder (Hall & Soskice, 2001, s. 30).

Organisering

Arbejdsmarkedene i LME'er er typisk fleksible, med lave barrierer for ansættelse og afskedigelse. Dette giver virksomhederne mulighed for hurtigt at tilpasse deres arbejdsstyrke til skiftende markedsforhold. Lønfastsættelsen sker ofte på individuelt niveau eller gennem decentraliserede forhandlinger, hvilket giver virksomhederne større fleksibilitet i forhold til lønomkostninger. Fagforeninger spiller normalt en mindre rolle i LME'er sammenlignet med CME'er (Hall & Soskice, 2001, s. 30-31).

Konkurrence

Virksomheder i LME'er konkurrerer primært på baggrund af priser og radikal innovation. Det betyder, at de konstant stræber efter at udvikle nye produkter og tjenester eller at finde nye måder at producere eksisterende produkter billigere på. Denne konkurrenceform tilskynder til en dynamisk og innovativ økonomi, hvor virksomheder er nødt til at være på forkant med den teknologiske udvikling for at overleve (Hall & Soskice, 2001, s. 30-31).

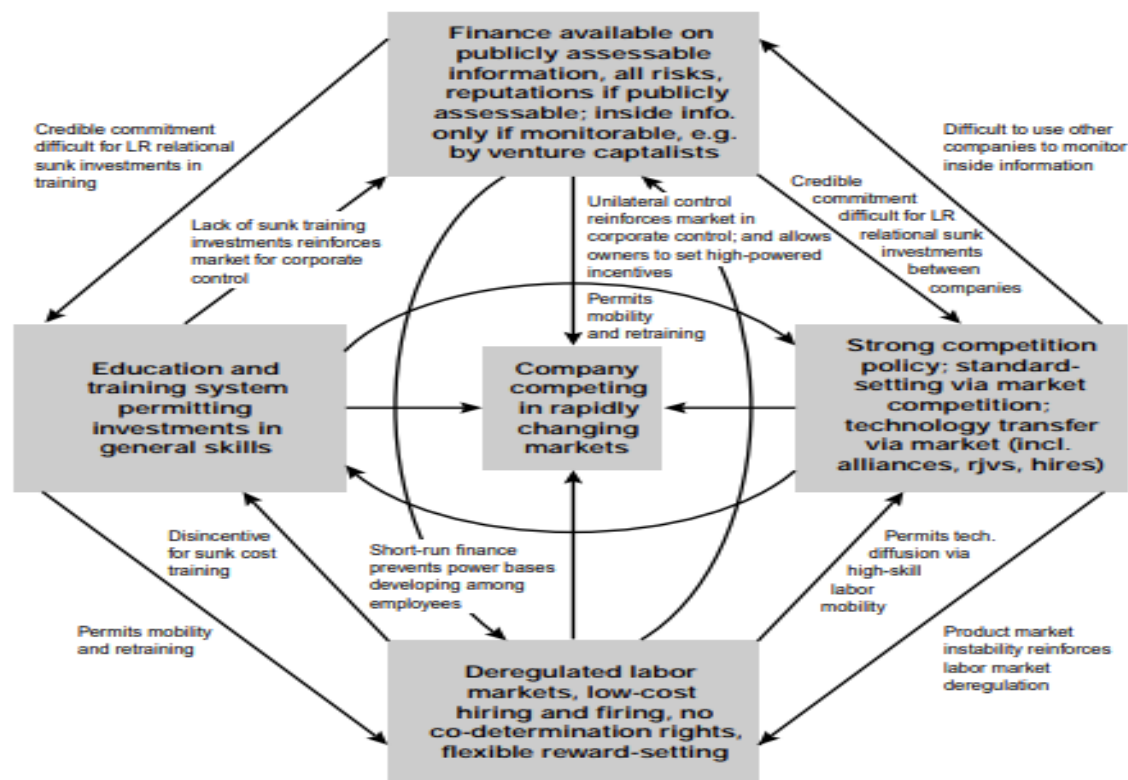


Fig. 1.4 Complementarities across subsystems in the American liberal market economy

Figur 1 Liberale markedsøkonomi (LME) | Hall & Soskice (2001)

Coordinated Market Economies (CME)

I koordinerede markedsøkonomier (CME'er) er økonomiske aktiviteter baseret på samarbejde, koordination og strategisk interaktion mellem aktører som virksomheder, fagforeninger og staten. CME'er, såsom Tyskland, er kendetegnet ved en høj grad af regulering, sociale partnerskaber og staten spiller en aktiv rolle i CME'er og er ofte involveret i at regulere markeder, støtte strategiske industrier og fremme social velfærd. I koordinerede markedsøkonomier er virksomheder mere afhængige af ikke-markedsbaserede relationer til at koordinere deres bestræbelser med andre aktører og opbygge deres kernekompetencer. Disse ikke-markedsbaserede former for koordination indebærer ofte mere omfattende relationelle eller ufuldstændige kontrakter, netværksovervågning baseret på udveksling af privat information inden for netværk, og mere tillid til samarbejdsorienterede, frem for konkurrenceprægede, relationer til at udvikle virksomhedens kompetencer. Dette kan skabe et mere stabilt og forudsigeligt erhvervsklima, der er gunstigt for langsigtede investeringer og samarbejde mellem virksomheder. Samtidig kan den statslige involvering også medføre bureaukratiske byrder og begrænse virksomhedernes fleksibilitet. I modsætning til liberale markedsøkonomier, hvor ligevægtsresultaterne af virksomhedernes adfærd normalt er bestemt af udbuds- og efterspørgselsforhold på konkurrencedygtige markeder, er ligevægtene, i CME'er som virksomhederne koordinerer i koordinerede, mere et output af strategiske interaktioner mellem virksomheder og andre aktører (Hall og Soskice, 2001, s.8.)

Finansiering

Virksomheder i CME'er har ofte adgang til langsigtet finansiering fra banker og andre finansielle institutioner, der er villige til at investere i projekter med et længere tidsperspektiv. Denne type finansiering er ikke afhængig af offentlige økonomiske data eller aktuelle afkast, men baserer sig i stedet på en vurdering af virksomhedens langsigtede potentiale. Investorer i tålmodig kapital accepterer lavere afkast på kort sigt for at sikre en stabil og sikker investering med højere afkast på lang sigt. For at sikre værdien af deres investeringer engagerer investorer sig i virksomhederne. Denne proces giver dem adgang til informationer om virksomhedens resultater, som normalt betragtes som private. Adgangen til disse informationer sker gennem den strategiske interaktion, hvor virksomheder deltager i tætte netværk og skaber nære relationer med hinanden, leverandører

og kunder. Gennem denne interaktion udvikles pålidelige oplysninger om virksomhedens fremskridt. Brancheorganisationer og fagforeninger inddrages også i samarbejdet og styrker pålideligheden af informationerne ved at kunne sanktionere virksomheder for vildledning (Hall & Soskice, 2001, s. 22-23).

Uddannelse

Arbejdsmarkedet i den koordinerede markedsøkonomi søger en arbejdsstyrke med høje virksomheds- eller industrispecifikke færdigheder. Dette skaber et koordineringsproblem, da arbejdsstyrken kun vil tilegne sig disse færdigheder, hvis det sikrer dem en lukrativ beskæftigelse. Samtidig skal virksomheder, der investerer i uddannelse, sikre sig, at arbejdsstyrken tilegner sig brugbare færdigheder og derefter bliver i virksomheden. Den strategiske interaktion løser dette koordineringsproblem gennem stærke arbejdsgiverforeninger og fagforeninger. Disse organisationer fører tilsyn med uddannelsessystemet og presser store virksomheder til at oprette lærepladser. De samarbejder også med virksomhederne om at definere de nødvendige kvalifikationskategorier og branchefærdigheder. Dette sikrer, at uddannelserne matcher virksomhedernes behov, og at der er en efterspørgsel efter færdiguddannede (Hall & Soskice, 2001, s. 25-26).

Organisering

Stærke brancheorganisationer, fagforeninger og arbejdsgiverforeninger spiller en afgørende rolle i den strategiske interaktion. De skaber samarbejde i hele systemet, fra sikring af dygtige erhvervsuddannelser til forhandling af langtidskontrakter og lønninger. De indsamler informationer om virksomhederne gennem koordinering af standardindstillinger og teknologioverførsel. Dette sikrer netværksovervågning, som muliggør finansiering af langsigtede afkast. Virksomhedsledelsen er ofte præget af konsensusbeslutninger og inddragelse af medarbejdere gennem medbestemmelse, hvilket kan føre til en mere ansvarlig og bæredygtig udvikling. Dette gør virksomhederne mindre udsatte for økonomiske nedture og giver dem mulighed for at ansætte arbejdsstyrken på lang sigt. I den koordinerede markedsøkonomi spiller organisationerne en afgørende rolle for at sikre, at systemet fungerer (Hall & Soskice, 2001, s. 26).

Konkurrence

Virksomheder i CME'er konkurrerer ofte på baggrund af kvalitet og inkrementel innovation, hvor de fokuserer på løbende forbedringer af eksisterende produkter og processer (Hall & Soskice, 2001, s. 16). Dette understøttes af et koordineret arbejdsmarked med stærke fagforeninger og arbejdsgiverorganisationer, der forhandler kollektive overenskomster. Disse overenskomster sikrer arbejdstagerne gode løn- og arbejdsvilkår, hvilket igen fremmer stabilitet og langsigtede investeringer i uddannelse og forskning (Hall & Soskice, 2001, s. 24-26). Virksomheder i CME'er er ofte villige til at investere i branchespecifikke færdigheder og samarbejde med andre virksomheder og uddannelsesinstitutioner for at sikre, at arbejdsstyrken har de nødvendige kompetencer (Hall & Soskice, 2001, s. 27).

28

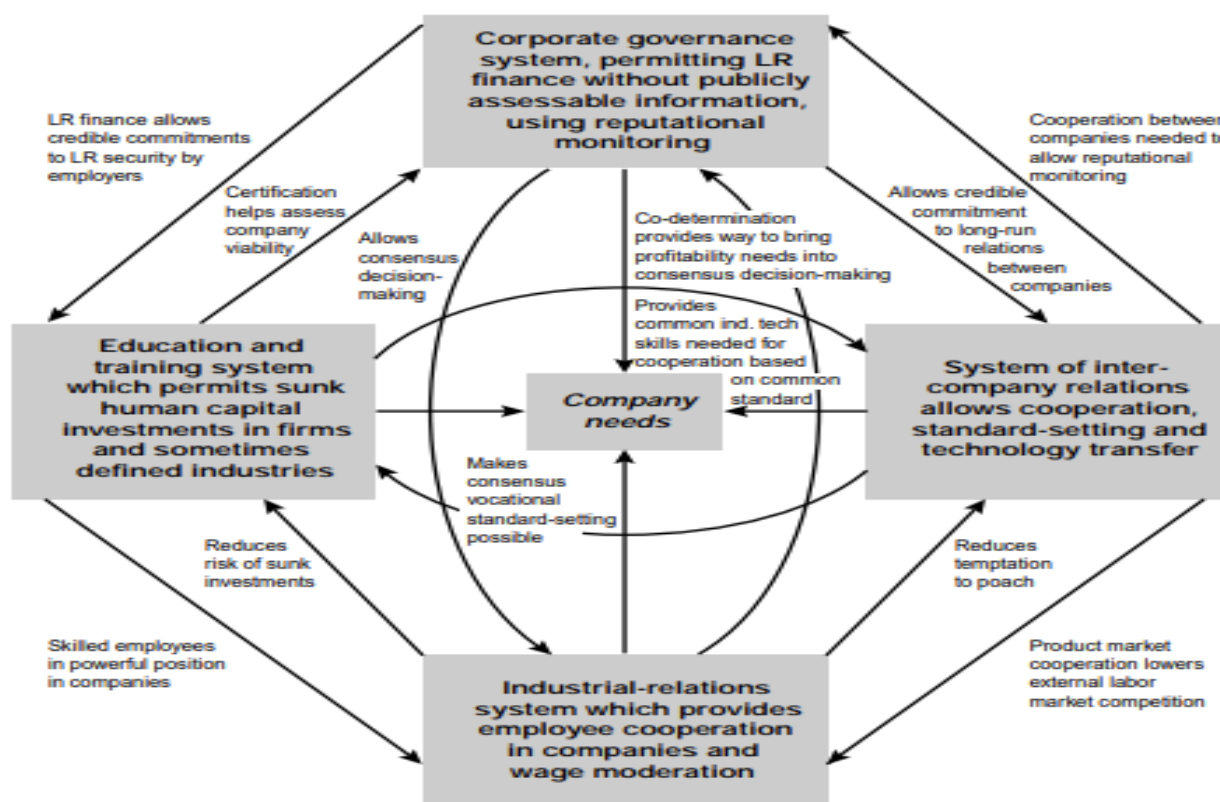
Peter A. Hall and David Soskice

Fig. 1.3 Complementarities across subsystems in the German coordinated market economy

Figur 2 Koordinerede markedsøkonomi (CME) | Hall & Soskice (2001)

Sammendrag af VoC idealtyper og institutionelle forskelle

Liberale markedsøkonomier (LME'er) og koordinerede markedsøkonomier (CME'er) udgør to distinkte idealtyper inden for varierede markedsøkonomier (VoC). De adskiller sig fundamentalt i institutionelle infrastrukturer, koordineringsmekanismer, innovative kapaciteter og dermed i makroøkonomiske resultater.

LME'er er præget af et armslængdeprincip og formelle juridiske kontrakter. Deres institutionelle rammer er mere fleksible og tilgodeser hurtige tilpasninger til skiftende markedsvilkår. I kontrast hertil har CME'er et tættere netværk af uformelle aftaler, netværk og strategisk koordinering. Deres institutionelle rammer er mere stabile og fokuserer på langsigtet vækst og stabilitet. Koordineringsmekanismerne adskiller sig også markant. I LME'er koordinerer virksomheder primært via markedet og formelle kontrakter. Konkurrencen er høj, og virksomheder fokuserer på profitabilitet og markedsandele. I CME'er er der en højere grad af samarbejde og koordinering mellem virksomheder, ofte via uformelle aftaler og netværk.

Disse forskelle i institutionelle infrastrukturer og koordineringsmekanismer fører til forskellige innovative kapaciteter. LME'er er typisk bedre til radikale innovationer, da deres rammer fremmer entreprenørskab og hurtige beslutningsprocesser. CME'er er derimod bedre til inkrementelle innovationer, da de fokuserer på langsigtede investeringer og specifikke fagkundskaber, der understøtter løbende forbedringer af eksisterende produkter og processer. De innovative kapaciteter afspejles i makroøkonomiske resultater. LME'er oplever typisk højere niveauer af entreprenørskab og dynamik, men kan også være mere volatile og præget af ulighed. CME'er har ofte en mere stabil vækst og lavere ulighed, men kan være mindre agile i forhold til hurtige markedsændringer.

LME'er og CME'er repræsenterer forskellige tilgange til markedsøkonomier. LME'er er mere fokuserede på markedskoordinering, fleksibilitet og radikal innovation, mens CME'er prioriterer koordinering, stabilitet og inkrementel innovation.

EU's Variety of Capitalism

EU består af 27 medlemslande med betydelige variationer mellem nord og syd samt landene fra øststudvidelsen. På trods af disse regionale forskelle, der er langt større end i Kina, har VoC-litteraturen i løbet af sin udvikling defineret en dominerende form for kapitalisme i EU. Bruno Amable (2003) identificerer fem hovedtyper af kapitalisme i Europa: den markedsbaserede model,

den socialdemokratiske model, den asiatiske model, den middelhavske model og den kontinentaleuropæiske model.

Ifølge Amable (2003) har EU-landene valgt eller fravalgt forskellige modeller afhængigt af deres politiske præferencer, sociale kompromiser og økonomiske udfordringer. For eksempel har de nordiske lande som Danmark, Sverige og Finland valgt den socialdemokratiske model, der kombinerer en høj grad af markedsfleksibilitet med en omfattende velfærdsstat og et stærkt uddannelsessystem. De sydeuropæiske lande som Spanien, Italien og Grækenland har derimod fravalgt denne model til fordel for den middelhavske model, der er kendetegnet ved en lav grad af markedsfleksibilitet, en svag velfærdsstat og et fragmenteret uddannelsessystem.

Den vigtigste type af kapitalisme i Europa er den koordinerede model, der er baseret på ideen om, at økonomisk succes afhænger af samarbejde og koordination mellem forskellige aktører som staten, arbejdsgiverne, fagforeningerne og civilsamfundet. Denne model er typisk for lande som Tyskland, Holland og Østrig, der har et højt niveau af markedsregulering, et korporatistisk arbejdsmarkedssystem, en generøs men selektiv velfærdsstat og et dualt uddannelsessystem med fokus på erhvervsuddannelse. Den koordinerede model kan også betragtes som en variant af den kontinentaleuropæiske model, der ifølge Amable (2003) er den mest udbredte form for kapitalisme i EU. Denne model er karakteriseret ved en moderat grad af markedsfleksibilitet, en blandet velfærdsstat med både universelle og målrettede ydelser, et hierarkisk uddannelsessystem med fokus på akademisk uddannelse og et diversificeret innovationssystem med både store virksomheder og små og mellemstore virksomheder.

Efter Brexit er det blevet mere erkendeligt, at EU i højere grad er en kontinental europæisk model end før. Det skyldes, at Storbritannien, som repræsenterede den liberale markedsmodel, har forladt EU. Storbritanniens udtræden betyder, at EU mistet en vigtig modvægt til lande som Frankrig, Italien og Tyskland, der nu i højere grad definerer EU's variety of capitalism. Det betyder dog ikke, at EU er en homogen blok af lande med den samme økonomiske model. Tværtimod er der stadig stor diversitet og variation mellem de forskellige EU-lande på tværs af Amables modeller (Rapacki & Czerniak, 2018). Der er også dynamik og forandring i de europæiske kapitalismer, som reagerer på nye udfordringer som globalisering, epidemier, krig, digitalisering, klimaforandringer og social ulighed. Derfor er det vigtigt at fortsætte med at analysere og forstå EU's variety of capitalism, både i teori og i praksis.

Kinas Variety of Capitalism

Den første udgave af VoC beskriver ikke den kinesiske form for kapitalisme, som for alvor tog fart efter årtusindskiftet. Rolf (2021) kritiserer begrænsningerne i Varieties of Capitalism i forståelsen af Kinas vækst og fremhæver problemer med institutioner, globale produktionsnetværk og intern geografisk udvikling. Herudover diskuteres manglerne ved heterodoks komparative kapitalisme. Rutten (2013) beskriver, at den kinesiske økonomi er kendetegnet ved en sameksistens mellem kapitalistiske institutioner og et autoritært system for økonomisk administration. Dette udfordrer traditionelle dikotomiske rammer for kapitalisme/socialisme eller koordinerede/markedsøkonomier. Det nuværende system i Kina er kendetegnet ved en markant industriel orientering, hvor stats- og markedsinstitutioner støtter hinanden. Det leninistiske apparat med bureaukratisk kontrol har skabt en dynamik, hvor økonomiske resultater fører til politisk indflydelse og kontrol med kapitalen. De finansielle markeder supplerer industriens krav om kapital, mens omorganiseringen af den offentlige sektor og velfærdssystemet har reduceret det finansielle pres på virksomhederne. Witt (2010) analyserede Kinas politiske økonomi ved hjælp af varianterne af kapitalistisk tilgang, og konkluderede, at Kina ligner en liberal markedsøkonomi (LME) i mange aspekter. Med hensyn til forbindelser mellem virksomheder udviser Kina karakteristika for en LME med kortsigtede relationer for markedsadgang frem for langsigtet teknologisk samarbejde.

Kinas statsejede virksomheder har let adgang til finansiering gennem statsejede banker, men bankerne har ikke meget kontrol over statsejede virksomheder.

Den statsejede sektor i Kina har ikke et bank- eller markedsbaseret system for virksomhedsledelse, hvilket har konsekvenser for interesserepræsentationen.

All-China Federation of Trade Unions (ACFTU) har en tendens til at gå på ledelsens side mod strejkende arbejdere og undlader at fremme arbejdernes interesser Rolf (2021).

China Enterprise Confederation (CEC) og All-China Federation of Industry and Commerce (ACFIC) er desuden statskontrollerede foreninger, hvor CEC har stærkere status hos statsejede virksomheder og ACFIC med private virksomheder. Ud fra litteraturen, så hælder Kinas variety of capitalism mod en koordineret markedsøkonomi, men med unikke karakteristika og kompleksitet i sine relationer mellem virksomheder og den statsejede sektor. Andet litteratur, forfattet af Zhang og Peck (2016), hævder, at den kinesiske kapitalisme model eksisterer i et trekantet forhold til den amerikanske liberale markedsøkonomi og den tyske koordinerede

markedsøkonomi. Litteraturen fremhæver også den betydelige grad af intern heterogenitet i det kinesiske tilfælde, med forskellige regionale stilarter af kapitalistisk udvikling, der forbliver forskellige fra hinanden. Specialets udgangspunkt er, at med ovenstående forbehold, at Kina er en særegen kombination af begge former idealtyper i form af LME og CME. EU har også en særdeles betydelige grad af intern heterogenitet af regionale variationer i form af Nordeuropa, Sydeuropa, Centraleuropa og Østeuropa.

Aktør	VoC	Tilgang
EU	Koordineret markedsøkonomi / (CME)	Fokus på konsensus og harmonisering af standarder. Forsigtig tilgang: Prioritering af databeskyttelse og etiske principper.
Kina	Statskapitalistisk markedsøkonomi	Fokus på national kontrol og teknologisk suverænitet. Topstyret regulering: Strenge databegrænsninger og fokus på national sikkerhed.
USA	Liberal markedsøkonomi (LME)	Fokus på innovation og minimal statslig indgriben. Tillid til selvregulering: Frivillige standarder og private initiativer.

Kritik af teori

VoC kan anskues som en approach og ikke en teori. Det er på den ene side en analytisk styrke, fordi VoC derved kan anlægges i forskellige varianter og dermed appliceres på forskelligartede empiriske problemstillinger. På den anden side sløres VoC's selvstændige videnskabelige bidrag,

hvis det blot betragtes som et perspektiv, der skal udfyldes teoretisk, inden det kan anvendes i empiriske analyser (Klitgaard og Roederer-Rynnin, 2009, s. 49).

Globaliseringsteori kunne belyse, hvordan den øgede globale integration påvirker nationale politiske økonomier og reguleringsmæssige rammer, mens teorier om interessevaretagelse kunne fremhæve, hvordan forskellige aktører påvirker den politiske beslutningsproces og dermed reguleringen af AI. VoC fokuserer primært på nationale institutionelle rammer og overser ofte den stigende betydning af globale aktører og processer i udformningen af AI-regulering. Tech-giganter som Google, Microsoft, Amazon og Meta har en enorm indflydelse på AI-landskabet på grund af deres store ressourcer, dataadgang og markedsdominans. Deres lobbyvirksomhed og strategiske partnerskaber med regeringer kan påvirke reguleringen i en retning, der gavner deres egne interesser, hvilket kan være i modstrid med de bredere samfundsmæssige interesser (Palley et al., 2023)

På trods af disse teories potentielle relevans, er de fravalgt i specialet for at opretholde et skarpt fokus på VoC-tilgangen. VoC giver en mere systematisk og sammenhængende ramme for at analysere de forskellige kapitalismemodeller og deres indflydelse på AI-reguleringen.

Specialet anerkender dog, at VoC-tilgangen har sine begrænsninger, især når det gælder forståelsen af AI som en ny og hastigt udviklende teknologi. VoC-tilgangen er primært udviklet til at analysere etablerede industrier og økonomiske sektorer, og den tager ikke fuldt ud højde for de unikke udfordringer og muligheder, som AI bringer med sig.

Derfor kunne fremtidige studier med fordel kombinere VoC-tilgangen med andre teoretiske perspektiver for at opnå en mere nuanceret og helhedsorienteret forståelse af AI-reguleringen. Ved at inddrage flere teorier og perspektiver kan der potentiel opnås bedre forståelse af de komplekse faktorer, der påvirker AI-reguleringen. Ydermere er videreudviklingen af teorien en svækkelse af forklaringskraften, da det i højere grad sætter sig i specialets kontekst og bliver sværere at kunne konkludere generelt ud fra teorien.

Metode

Dette afsnit præsenterer specialets metodiske tilgang, herunder en redegørelse for den videnskabsteoretiske position, den valgte teori og de anvendte metoder. Derudover diskuteres styrker og svagheder ved det metodiske valg og præsenteres en oversigt over de empiriske data, der anvendes i specialet. Afsnittet afsluttes med en præsentation af de undersøgelsesspørgsmål, der guider analysen.

Kvalitativ metode

Den kvalitative metode er velegnet til at undersøge komplekse fænomener som AI-regulering, da den giver mulighed for at indhente detaljerede og kontekstafhængige data. Dette er afgørende for at forstå de forskellige aktørers perspektiver og oplevelser, som spiller en central rolle i reguleringsprocessen. Et multipelt casestudie design tillader en komparativ analyse af AI-regulering i tre forskellige kontekster. Dette gør det muligt at identificere ligheder og forskelle i de forskellige reguleringsmæssige rammer og at udpege faktorer, der kan have betydning for effektiviteten af reguleringen.

Specialets metodevalg har en række styrker. Den kvalitative metode giver mulighed for en dybdegående forståelse af de forskellige aktørers perspektiver og oplevelser af AI-regulering. Det multiple casestudie design giver mulighed for at sammenligne og kontrastere reguleringen i tre forskellige landekontekster, hvilket bidrager til en mere nuanceret forståelse af emnet. Der er dog også svagheder forbundet med specialets metodevalg. Den kvalitative metode er mindre velegnet til at generalisere resultaterne til en bredere population. Desuden kan den kvalitative metode være vanskeligere at replicere end en kvantitativ metode.

Komparativt casedesign

Et casestudie er en forskningsstrategi, der beror på en detaljeret gennemgang af et konkret fænomen for at udlede, hvilken lære der kan uddrages fra netop dette fænomen. Gennem et komparativt casestudie-design undersøges AI-reguleringen i EU, USA og Kina. Disse cases er valgt, fordi de repræsenterer forskellige kapitalismemodeller (henholdsvis CME, LME og statskapitalisme) og har en betydelig indflydelse på den globale AI-udvikling. Casestudierne giver

mulighed for at identificere ligheder og forskelle i de forskellige reguleringsmæssige rammer og at undersøge, hvordan de institutionelle kontekster påvirker reguleringen af AI i praksis.

Specialet støtter sig til Robert Yins definition, der defineres casestudiet således:

“An empirical inquiry about a contemporary phenomenon (e.g. a “case”), set within its real-world context – especially when the boundaries between phenomenon and context are not clearly evident - (Yin, 2011, s. 4)”

Et komparativt casedesign involverer en dybdegående analyse af to eller flere cases. Ved at sammenligne disse cases kan man identificere mønstre og ligheder, der kan føre til konklusioner om det undersøgte fænomen (Bryman, 2016, s. 60-64). Dette speciale tager udgangspunkt i Varieties of Capitalism (VoC) teorien og anvender denne til at analysere, hvordan institutionelle kontekster i tre forskellige lande kan forme AI-regulering. Specialet har en komparativ tilgang, hvor AI-regulering i de tre cases sammenlignes og kontrasteres for at identificere ligheder, forskelle og mønstre. Formålet er at forklare, hvorfor forskellige varianter af kapitalisme kan føre til forskellige reguleringsmæssige rammer for AI. Specialet arbejder deduktivt, idet det starter med VoC-teorien, og bruger den til at analysere de tre cases. Den kritiske realisme danner grundlaget for specialet, med fokus på kritisk vurdering. Målet er at opnå en vis grad af generaliserbarhed ved at sammenligne de tre cases og identificere generelle mønstre og tendenser (Bryman, 2016, s. 64-70). Specialets designvalg, beror på argumentet, at enkelt casedesign oftest ikke er så overbevisende som et multipelt casestudie, hvis formålet er at teste eksisterende teori (de Vaus, 2011, s. 227).

Dokumentanalyse

Specialet anvender dokumentanalyse som sin primære dataindsamlingsmetode. Dette indebærer en systematisk gennemgang og analyse af relevante dokumenter, såsom lovforslag, bekendtgørelser, rapporter og politiske udtalelser. Dokumentanalysen giver adgang til de officielle holdninger og politikker for AI-regulering i de tre cases og muliggør en sammenligning af deres forskellige tilgange. En dokumentanalyse er en metode til at undersøge og fortolke skriftlige kilder, som kan være både primære og sekundære. Dokumentanalyse kan bruges til at belyse forskellige aspekter af et emne, såsom historie, kultur, politik, ideologi, værdier og normer. Dokumentanalyse kan også kombineres med andre metoder, som interviews eller observationer, for at opnå en dybere

forståelse af et fænomen. Dokumentanalyse er en nyttig metode til at udforske forskellige aspekter af et emne gennem skriftlige kilder. Dokumentanalyse kan give indsigt i historiske eller nutidige fænomener, samt udfordre eller bekræfte eksisterende antagelser eller teorier. Dokumentanalyse kræver dog en kritisk og reflektiv tilgang til valg, analyse og rapportering af dokumenter. Dokumenter udgør både empirien og analysens objekt i dokumentanalyse, som ofte bruges til at undersøge de skriftlige beslutninger, retningslinjer, processer eller overordnede mål indenfor et bestemt felt. Et vigtigt aspekt af dokumentanalyse er udvælgelsen af dokumenter til forskning og analyse: hvilke kriterier og begrundelser ligger bag valget af dokumenter (Lynggaard, 2010, s. 137). Der er tre typer af dokumenter: primære, sekundære og tertiære. Det er vigtigt at skelne mellem disse typer, fordi de har forskellige tidspunkter for deres skabelse. De primære og sekundære dokumenter er dokumenter, der er skrevet i nær relation til den praksis eller proces, de handler om, mens de tertiære dokumenter er skrevet senere (Lynggaard, 2010, s. 138). Specialet har valgt primære dokumenter til dokumentanalysen. Dokumentanalyse som metode har en række styrker, herunder muligheden for at få indblik i officielle holdninger og politikker, samt at kunne sammenligne og kontrastere forskellige reguleringsrammer. Dog er der også svagheder, såsom risikoen for bias i udvælgelsen og fortolkningen af dokumenter, samt den begrænsede mulighed for at undersøge årsagssammenhænge og inddrage aktørernes perspektiver (Lynggaard, 2010).

Dataindsamling

I dette speciale benyttes en komparativ dokumentanalyse til at undersøge, hvordan forskellige kapitalismemodeller påvirker reguleringen af kunstig intelligens (AI) i EU, USA og Kina. Valget af dokumenter er faldet på EU's AI Act, USA's bekendtgørelse om AI, og Kinas 24 retningslinjer for kunstig intelligens. Disse dokumenter er udvalgt, fordi de repræsenterer de officielle holdninger og politikker for AI-regulering i de tre cases, og dermed giver et indblik i, hvordan de forskellige samfundsmodeller og institutionelle rammer påvirker deres tilgang til AI.

For at sikre en pålidelig og troværdig analyse er det afgørende at vurdere dokumenternes kvalitet. I dette speciale er der anvendt J. Scotts (1990) fire kriterier for dokumentanalyse til at evaluere de udvalgte dokumenter:

Autenticitet: Dokumenterne vurderes som autentiske, da de er officielt udstedt af de relevante myndigheder i EU, USA og Kina. Dette sikrer, at dokumenterne er, hvad de giver sig ud for at

være, og at de er forfattet af de instanser, der har autoritet til at lovgive om AI. Der er ingen tegn på forfalskning eller manipulation, da de er offentligt tilgængelige.

Troværdighed: Dokumenterne vurderes som troværdige, da de er baseret på forskellige kilder, herunder forskning, policy-dokumenter og offentlige høringer. Dette indikerer, at de er resultatet af en grundig og velinformeret proces. Derudover er dokumenterne nøjagtige og pålidelige, da de afspejler de officielle holdninger og politikker for AI i de tre cases.

Repræsentativitet: Selvom dokumenterne repræsenterer de officielle holdninger til AI-regulering i hver region, er det vigtigt at anerkende, at de ikke nødvendigvis repræsenterer alle interessenters synspunkter. Der kan være divergerende holdninger og perspektiver, som ikke kommer til udtryk i disse dokumenter. Dette er en begrænsning, der bør adresseres i specialets diskussion og konklusion.

Mening: Dokumenternes formål er klart defineret. De søger at etablere juridiske rammer og principper for AI i overensstemmelse med de grundlæggende værdier hver case har. Dog kan der være forskellige fortolkninger af, hvordan disse mål skal opnås, og hvilke værdier der skal prioriteres. Dette kan føre til forskellige implementeringer og konsekvenser af AI-reguleringen i praksis.

For at imødekomme den potentielle bias eller skjulte dagsordener i dokumenterne, vil specialet anvende en kritisk tilgang til analysen. Ved at være opmærksom på de forskellige aktørers interesser og magtforhold, samt de historiske og institutionelle kontekster, hvori dokumenterne er blevet til, kan specialet bidrage til en mere nuanceret og kritisk forståelse af, hvordan AI-regulering formes og implementeres i forskellige kapitalistiske modeller.

Forskningskriterier

Kvalitativ forskning adskiller sig fra kvantitativ forskning ved ikke at have de samme standardiserede kriterier for evaluering af forskningens kvalitet. Hvor kvantitativ forskning ofte læner sig op ad målbare kriterier som reliabilitet (konsistens i målingerne) og validitet (gyldigheden af resultaterne), kræver kvalitativ forskning en mere nuanceret tilgang, der tager højde for forskningens kontekst, formål og metodiske valg (Bryman, 2016, s. 41).

Intern validitet, der i kvantitativ forskning handler om at sikre en gyldig kausalslutning (Bryman, 2016, s. 41), kan i kvalitativ forskning oversættes til troværdighed. Det handler om, hvorvidt undersøgelsens resultater giver et troværdigt og overbevisende billede af den undersøgte

virkelighed. I dette speciale er der sikret intern validitet ved grundigt at redegøre for de teoretiske begreber og operationalisere dem, hvilket vil sige at tydeliggøre, hvordan de anvendes i praksis. Derudover er der anvendt forskellige metoder til at indsamle data, hvilket styrker validiteten gennem triangulering (Bryman, 2016, s. 386).

For at sikre målingsvaliditet i specialet er der anvendt en kombination af teoretiske og operationelle definitioner af de centrale begreber (Bryman, 2016, s. 151). De teoretiske definitioner, der er baseret på VoC-teorien og historisk institutionalisme, beskriver, hvad begreberne betyder i en bredere kontekst, mens de operationelle definitioner tydeliggør, hvordan begreberne konkret anvendes i analysen af de empiriske data. Denne fremgangsmåde styrker specialets målingsvaliditet, da det sikrer en klar forbindelse mellem teori og empiri. Derudover kan det overvejes at inddrage respondentvalidering (Bryman, 2016, s. 385), hvor respondenterne får mulighed for at kommentere på resultaterne, som en yderligere sikring af målingsvaliditeten. Kvalitativ forskning og dens kvalitet skal vurderes ud fra kriterier, der anerkender forskningens mål og ambitioner (Brinkmann & Tanggaard, 2020, s. 659). Validitet indenfor kvalitative metoder, har mere fokus på de sproglige ytringer. Sproglige ytringer skal afspejle det for undersøgelsens relevante fagområder, og kommer til at omhandle refleksion over betydning, underbetydning og merbetydning (Ingemann et al., 2018, s. 60). Kvalitativ forskning spiller sproget en central rolle, da det er gennem sproget, at vi forstår og fortolker den sociale verden. I specialet er der lagt vægt på at analysere de sproglige ytringer i de forskellige dokumenter. Ved at være opmærksom på de sproglige nuancer kan specialet opnå en dybere forståelse af de forskellige aktørers perspektiver og interesser.

Styrker og svagheder ved metodevalget

Styrker	Svagheder
Dybdegående forståelse af aktørers perspektiver og oplevelser.	Begrænset generaliserbarhed af resultaterne.
Mulighed for at undersøge komplekse fænomener i deres kontekst.	Udfordringer med replikation af studiet på grund af den kvalitative metodes subjektive karakter.
Adgang til officielle holdninger og politikker gennem dokumentanalyse.	Risiko for bias i udvælgelsen og fortolkningen af dokumenter.
Mulighed for at sammenligne og kontrastere forskellige reguleringsrammer gennem et multipelt casestudie-design.	Begrænset mulighed for at undersøge årsagssammenhænge mellem institutionelle rammer og AI-regulering, da casestudier ikke er velegnede til at isolere effekten af enkelte variable.
Kritisk realismes fokus på at afdække underliggende mekanismer og strukturer.	Risiko for at overse vigtige aktører og processer, der ikke er repræsenteret i de udvalgte dokumenter.

Empiri

Specialet undersøger reguleringen af kunstig intelligens (AI) i tre forskellige kontekster: EU, USA og Kina. Disse cases er valgt strategisk, da de repræsenterer forskellige kapitalismemodeller (henholdsvis CME, LME og statskapitalisme) og har en betydelig indflydelse på den globale AI-udvikling. For EU er AI Act valgt som det primære dokument, da det er den mest omfattende og ambitiøse AI-regulering på globalt plan. For USA er præsident Bidens dekret valgt, da den repræsenterer USA's nyeste og mest omfattende holdning til AI-regulering. For Kina er de 24 retningslinjer for kunstig intelligens, da reguleringen er nyere og i højere grad giver indsigt i Kinas tilgang til AI-regulering efter ChatGPT modsat tidligere regulering.

Den komparative analyse fokuserer udelukkende på disse tre aktører, hvilket kan medføre en begrænset forståelse af de globale konsekvenser af AI-regulering. Selvom de tre cases har gjort betydelige fremskridt inden for AI, befinder de sig på forskellige stadier af implementeringen af deres AI-regulering.

Selvom de udvalgte styringsdokumenter opfylder de samme kriterier for relevans, er deres beskaffenhed ikke ensartet. For eksempel er EU's AI Act et omfattende og gennearbejdet lovforslag med detaljerede bestemmelser, mens USA's dekret er en mere overordnet retningslinje udstedt som dekret udenom Kongressen.

En central afgrænsning er, at analysen fokuserer på de tre udvalgte styringsdokumenter og ikke inddrager de mange andre AI-relaterede politikker og initiativer, der findes i de tre case. For eksempel har EU vedtaget en række andre digitale reguleringer, såsom GDPR, Digital Markets Act og Digital Services Act, der også har betydning for AI-reguleringen, da de regulerer algoritmer og tech-giganterne. Ligeledes har Kina implementeret en lang række sektorspecifikke regler og retningslinjer for AI. Denne afgrænsning er nødvendig for at sikre en fokuseret analyse, men det betyder også, at specialet ikke giver et fuldstændigt billede af AI-reguleringen i de tre cases.

Endelig afgrænser specialet sig fra at undersøge magtdynamikken mellem de tre cases, selvom dette spiller en afgørende rolle i deres AI-strategier og -reguleringer. Denne afgrænsning er foretaget på grund af specialets begrænsede omfang, men det skal anerkendes, at den geopolitiske konkurrence mellem USA og Kina og EU's rolle som en potentiel reguleringsmæssig leder kan have betydelige konsekvenser for den fremtidige udvikling af AI-regulering på globalt plan. For eksempel kan EU's AI Act potentielt påvirke andre lande gennem den såkaldte "Bruselles effect", hvor EU's regulering bliver en de facto global standard ved at inspirere andre regeringer.

På trods af disse afgrænsninger giver specialet et værdifuldt bidrag til forståelsen af, hvordan forskellige kapitalismemodeller påvirker reguleringen af AI. Ved at fokusere på EU, USA og Kina som cases kan specialet belyse de centrale forskelle og ligheder i deres tilgange og dermed bidrage til den bredere debat om, hvordan AI bedst kan reguleres i en globaliseret verden.

EU's AI Act

I 2021 præsenterede EU-Kommissionen et forslag til en forordning om kunstig intelligens (AI Act), der sigter mod at skabe et juridisk rammeværk for AI-udvikling og -anvendelse i EU. Forslaget kom i en tid med stigende bekymringer om de potentielle risici og etiske dilemmaer forbundet med AI, men også med en erkendelse af de potentielle fordele for samfundet og den europæiske økonomi. Kommissionens forslag fra 2021 var et af de første forsøg på at skabe et omfattende juridisk rammeværk for AI på globalt plan. Det har haft stor indflydelse på den efterfølgende debat om AI-regulering i EU og globalt.

Rådet færdiggjorde sine forslag til ændringer af AI Act i slutningen af 2022, hvorefter Europa-Parlamentet (EP) fremlagde sin version i sommeren 2023 (European Commission, 2023). Få af de foreslåede ændringer berørte dog direkte eller indirekte spørgsmålet om AI-suverænitet. I præamblerne til lovgivningen understregede EP de miljømæssige og samfundsmæssige konsekvenser af AI-systemer (European Commission, 2023, s. 6). Selvom denne betoning i lovgivningen er vag, er den mere fremtrædende i den endelige version af AI Act end i Kommissionens oprindelige forslag. EP's ændringsforslag understregede også, at AI-systemer ikke bør underminere borgernes autonomi (European Parliament et al., 2022, s. 13), og udtrykte bekymring for, at fordelene ved AI og investeringer i dem vil blive fordelt ujævnt i hele unionen (European Parliament et al., 2022, s. 17). Samtidig foreslog man at tilføje sprog til Kommissionens udkast, der understreger betydningen af "AI fremstillet i Europa" (European Parliament et al., 2022, s. 14).

Set i lyset af de samlede trilogforhandlinger mellem EP, Rådet og Kommissionen, har forhandlingerne ikke grundlæggende ændret AI Act's tilgang til suverænitetsspørgsmålet. Frankrig, Tyskland og Italiens regeringer kæmpede mod en for streng regulering, da landene mente, at det kunne hæmme konkurrenceevnen for europæiske virksomheder. Med undtagelse af bestemmelserne om generative AI-systemer, der blev tilføjet undervejs efter ChatGPT udkom, er den endelige lovgivningstekst i høj grad Kommissionens oprindelige forslag fra 2021.

USA's executive order om kunstig intelligens

Den 30. oktober 2023 udstedte præsident Joe Biden en Executive Order om sikker, sikker og pålidelig udvikling og brug af kunstig intelligens (AI). Dette skridt markerer en betydningsfuld milepæl i USA's tilgang til regulering af AI-teknologier. Baggrunden for denne beslutning er mangefacetteret og afspejler både nationale og internationale bekymringer samt muligheder.

For det første anerkender ordren AI's ekstraordinære potentiale til at løse presserende udfordringer og bidrage til velstand, produktivitet, innovation og sikkerhed. Samtidig fremhæver den de risici, der er forbundet med uansvarlig brug af AI, såsom svindel, diskrimination, bias og misinformation. Målgruppen er bred og omfatter føderale myndigheder, AI-udviklere og offentligheden. For føderale myndigheder definerer dokumentet direktiver for udvikling og brug af AI-systemer. AI-udviklere præsenteres for forventninger til at sikre, at deres systemer er sikre, pålidelige og etiske.

Endelig informerer det offentligheden om regeringens vision for en ansvarlig og gavnlig anvendelse af AI.

Der identificeres fem grundlæggende retningslinjer, som skal styre udviklingen og implementeringen af kunstig intelligens. Disse inkluderer beskyttelse af persondata, fremme af lighed, modvirkning af forudindtagethed, garanti for systemernes sikkerhed og robusthed, samt sikring af ansvarlighed. Føderale institutioner pålægges at implementere nødvendige foranstaltninger for at sikre, at AI-systemer er beskyttede og har modstandsdygtighed overfor cybertrusler. Der ydes støtte til forskning og udvikling af AI-teknologier, der er i overensstemmelse med de etablerede retningslinjer. Regeringen forpligter sig til internationalt samarbejde for at sikre en ansvarsfuld og fordelagtig brug af kunstig intelligens.

Joe Bidens executive order er et væsentligt dokument, der giver indblik i den amerikanske regerings politik for AI. Den præsenterer en række principper og retningslinjer, der har til formål at styre udviklingen og brugen af AI på en ansvarlig og gavnlig måde. Dokumentet er relevant for en bred målgruppe og kan tjene som et nyttigt udgangspunkt for videre diskussion og analyse af AI-teknologiens fremtid. Udviklingen af ordren kan ses som en direkte respons til lignende lovgivningsmæssige tiltag fra andre globale magter såsom EU og Kina, som begge har taget skridt til at forme fremtiden for AI gennem deres egne reguleringsrammer. EU har været særligt proaktiv med deres hvidbog om kunstig intelligens, som lægger op til strenge krav til AI-systemers gennemsigtighed og ansvarlighed.

Kinas 24 retningslinjer for kunstig intelligens

I de seneste to år har Kina implementeret nogle af verdens første bindende nationale regulativer for kunstig intelligens (AI). Disse regulativer omfatter blandt andet anbefalingsalgoritmer til indholdsspredning, syntetisk genererede billeder og videoer, samt generative AI-systemer som OpenAI's ChatGPT. Reglerne indfører nye krav til konstruktion og implementering af algoritmer, og til hvilke oplysninger AI-udviklere skal afsløre for regeringen og offentligheden. Disse tiltag skaber det intellektuelle og bureaukratiske fundament for en omfattende national AI-lov, som Kina sandsynligvis vil offentliggøre i de kommende år. Dette udgør en potentielt betydningsfuld udvikling for global AI-styring på niveau med Den Europæiske Unions kommende AI-lov. Samlet

set omdanner disse bevægelser Kina til et laboratorium for eksperimenter i styringen af måske denne æras mest indflydelsesrige teknologi.

Disse regulativer er ikke kun relevante for Kina, men også for resten af verden, da de kan sætte en standard for global AI-regulering. Med Kinas hurtige udvikling inden for AI og dets ambitioner om at blive en førende magt på området, er det afgørende at forstå landets tilgang til regulering af denne transformative teknologi. Det er en balancegang mellem at fremme innovation og sikre, at teknologien ikke misbruges eller skaber uforudsete negative konsekvenser.

I april 2023 udsendte regulatorerne et udkast til den nye regulering for generative AI til offentlig høring. Dokumentet er oprindeligt udgivet af Cyberspace Administration of China (CAC), som er en officiel regeringsinstitution i Kina. Dette giver dokumentet en primær kildestatus, da det kommer direkte fra en regulerende myndighed med ansvar for cyberspacepolitikker i Kina. Dette dokument, udarbejdet af Cyberspace Administration of China (CAC), har til formål at etablere et reguleringsmæssigt rammeværk for Generative AI (GAI) i Kina. Det er oversat fra kinesisk til engelsk af Kina-eksperter på Stanford University. "Foranstaltningerne" adresserer flere centrale områder, herunder definitionen af GAI, kategorisering af GAI-tjenester, licenskrav for udbydere, sikkerheds- og etisk vurdering af GAI-modeller, databeskyttelse og privatliv samt tilsyn og håndhævelse. Dokumentet er opdelt i syv kapitler, der dækker alt fra generelle bestemmelser til tillægsbestemmelser. Det definerer GAI, kategoriserer tjenester baseret på risiko, fastlægger licenskrav, beskriver sikkerhedsvurderinger, adresserer databeskyttelse og privatliv, og definerer tilsynsmyndighedernes rolle.

Hypoteser

I dette afsnit fremsættes en række hypoteser, der vil blive undersøgt i specialet. Hypoteserne er udformet på baggrund af den teoretiske ramme og de problemstillinger, der er blevet præsenteret tidligere i specialet. Hypoteserne vil blive testet gennem en komparativ analyse af EU, USA og Kinas tilgang til regulering af kunstig intelligens (AI).

Hypotese 1

EU's AI-regulering vil afspejle en koordineret markedsøkonomi (CME) med vægt på sociale hensyn og forsigtighedsprincippet, men med en vis grad af fleksibilitet og tilpasningsevne som følge af de løbende diskussioner og forhandlinger mellem forskellige aktører i den europæiske kontekst.

USA's AI-regulering vil afspejle en liberal markedsøkonomi (LME) med fokus på innovation og mindre statslig indblanding, men med potentiale for øget statslig involvering, hvis der opstår betydelige markedsfejl eller trusler mod samfundsmæssige interesser.

Kinas AI-regulering vil afspejle en statskapitalistisk model med vægt på national sikkerhed og social stabilitet, men med en vis grad af åbenhed over for markedsbaseret innovation og internationalt samarbejde for at fremme økonomisk vækst.

Hypotese 2

EU's institutionelle rammer vil muliggøre en mere omfattende og inkluderende AI-regulering, der inddrager forskellige interessenter, men implementeringen af reguleringen vil være præget af kompromisser og forhandlinger mellem medlemslandene, hvilket kan føre til forsinkelser og udvanding af visse bestemmelser.

USA's institutionelle rammer vil begrænse omfanget af statslig indblanding i AI-regulering, men den politiske proces vil være præget af lobbyisme og interessekonflikter mellem forskellige aktører, herunder teknologivirksomheder, NGO'er og politiske partier.

Kinas institutionelle rammer vil give staten mulighed for at implementere en hurtig og effektiv AI-regulering, men den manglende inddragelse af civilsamfundet og uafhængige eksperter kan øge risikoen for utilsigtede konsekvenser og begrænse teknologiens potentiale.

Hypotese 3

EU's reguleringsramme vil være mere præventiv og risikobaseret end USA's, men den vil også indeholde fleksible elementer, såsom mulighed for regulatoriske sandkasser og undtagelser for forskning og innovation, for at fremme teknologiudvikling og undgå at hæmme økonomisk vækst.

USA's reguleringsramme vil være mere reaktiv og baseret på frivillige retningslinjer og standarder, men der vil være potentiale for en mere proaktiv tilgang, hvis der opstår alvorlige sikkerheds- eller etiske problemer i forbindelse med AI-anvendelser.

Kinas reguleringsramme vil være en blanding af statslig kontrol og støtte til udvalgte industrier, men der vil også være en vis grad af fleksibilitet og eksperimentering, især inden for særlige økonomiske zoner og innovationsklynger.

Hypotese 4

EU's tilgang til AI-regulering vil påvirke innovationens retning ved at skabe incitamenter for udvikling af AI-systemer, der er i overensstemmelse med europæiske værdier og standarder for etik, sikkerhed og gennemsigtighed.

USA's tilgang vil sandsynligvis betyde hurtigere udvikling og implementering af AI-teknologier, men også til en større risiko for utilsigtede konsekvenser og ulige fordeling af fordele og ulemper. Kinas tilgang vil give staten mulighed for at målrette investeringer og ressourcer mod strategisk vigtige AI-områder, men kan også begrænse den overordnede innovation og skabe etiske dilemmaer i forhold til overvågning og kontrol.

Undersøgelsesspørgsmål

Jævnfør problemformuleringen, den tilgrundliggende teori og ovenstående hypoteser udformes følgende undersøgelsesspørgsmål, som sætter retningen for analysen:

1. Hvordan afspejler regulering af kunstig intelligens de respektive aktørers varieties of capitalism?

- Hvilke værdier og principper er fremherskende i de forskellige reguleringsrammer?
- Hvordan afspejler reguleringen de forskellige økonomiske modeller og institutionelle strukturer?
- Hvilke aktører har haft størst indflydelse på udformningen af reguleringen?

Dette spørgsmål og underspørgsmål undersøger sammenhængen mellem den specifikke udformning af AI-regulering og den type kapitalisme, der er fremherskende i et land. For eksempel kan man forvente, at regulering i et land med koordineret markedsøkonomi (CME) som EU vil lægge mere vægt på social ansvarlighed og forbrugerbeskyttelse, sammenlignet med et land med liberal markedsøkonomi (LME) som USA, hvor fokus i højere grad er på innovation og fri konkurrence.

2. Hvordan påvirker de institutionelle kontekster i koordinerede markedsøkonomier (EU), statskapitalisme (Kina) og liberale markedsøkonomier (USA) tilgangen til AI-regulering?

- Hvilke institutioner og aktører er involveret i AI-reguleringen i de forskellige cases?
- Hvordan påvirker de forskellige institutionelle rammer beslutningsprocessen og implementeringen af AI-regulering?
- Hvilke typer af koordinering og samarbejde finder sted mellem aktørerne i de forskellige cases?

De institutionelle rammer, der er til stede i et land, spiller en afgørende rolle for, hvordan AI-regulering udformes og implementeres. Spørgsmålet og underspørgsmål undersøger derfor, hvordan de forskellige institutionelle kontekster i CME, statskapitalisme og LME påvirker tilgangen til AI-regulering. For eksempel kan man forvente, at et land med et stærkt fokus på planlægning og koordinering som Kina vil have en mere proaktiv tilgang til AI-regulering, sammenlignet med et land med en laissez-faire tilgang som USA.

3. Hvad er de specifikke reguleringsrammer for AI-udvikling?

- Hvilke love, politikker, retningslinjer og standarder er relevante for AI-udvikling og -anvendelse?
- Hvordan defineres og klassificeres AI-systemer i de forskellige reguleringsrammer?
- Hvilke krav og forpligtelser stilles der til udviklere og brugere af AI-systemer?

For at sammenligne og analysere de forskellige tilgange til AI-regulering er det nødvendigt at undersøge de specifikke reguleringsrammer, der er på plads i de tre udvalgte cases. Dette inkluderer en analyse af lovgivning, politikker og retningslinjer, der er relevante for AI-udvikling og brug.

Analyse

I analysen præsenteres en dybdegående undersøgelse af, hvordan reguleringen af kunstig intelligens (AI) afspejler de forskellige kapitalistiske systemer, som aktørerne opererer indenfor. Med udgangspunkt i en solid teoretisk ramme, adresserer analysen først spørgsmålet om, hvordan AI-reguleringen spejler den kapitalistiske variation mellem lande. Her er der en antagelse om, at

lande med koordinerede markedsøkonomier som EU, vil have en tendens til at regulere AI med større vægt på social ansvarlighed og forbrugerbeskyttelse, i modsætning til liberale markedsøkonomier som USA, hvor innovation og konkurrence prioriteres højere. Dernæst udforsker analysen, hvordan de institutionelle kontekster i forskellige økonomiske systemer - koordinerede markedsøkonomier, statskapitalisme og liberale markedsøkonomier - former tilgangen til AI-regulering. Dette inkluderer en vurdering af, hvordan lande som Kina, med en statskapitalistisk model, muligvis vil have en mere proaktiv og planlagt tilgang til AI-regulering sammenlignet med mere laissez-faire orienterede lande som USA. Analysen bevæger sig videre til at sammenligne de specifikke reguleringsrammer for AI, som er etableret i de udvalgte cases, og undersøger lovgivning, politikker og retningslinjer, der er relevante for udviklingen og brugen af AI.

EU

EU's AI-lovgivning er designet til at sikre, at AI-systemer respekterer grundlæggende rettigheder, sikkerhed og etiske principper. Dette afspejler en model, hvor markedet er koordineret gennem et netværk af relationer og forpligtelser, der går ud over ren markedslogik. I koordinerede markedsøkonomier, er der en tendens til at regulere AI på en måde, der fremmer samarbejde mellem staten, virksomheder og arbejdsmarkedets parter. Komplexiteten i lovgivningen afspejler de forskellige nuancer af overvejende koordinerede modeller, der findes inden for EU. Denne analyse undersøger forholdet mellem EU's AI-lovgivning og varieties of capitalism-teorien.

Årsagerne til EU's ønske om at regulere AI er mangesidede og afspejler både de enorme muligheder og potentielle risici, som AI-teknologier indebærer. På den ene side har AI potentialet til at revolutionere en lang række sektorer, fra sundhed og miljø til finans og offentlig administration, ved at forbedre forudsigelser, optimere operationer og ressourceallokering og personalisere tjenester. På den anden side rejser brugen af kunstig intelligens alvorlige spørgsmål om beskyttelsen af grundlæggende rettigheder og brugernes sikkerhed, når teknologien integreres i produkter og tjenester.

EU's AI-lov sigter mod at finde en balance mellem at fremme innovation og teknologisk udvikling og samtidig sikre, at AI bruges på en måde, der respekterer EU's værdier og standarder. Loven indfører en risikobaseret tilgang, hvor AI-systemer klassificeres efter deres potentielle risiko, og der stilles forskellige krav og forpligtelser baseret på denne klassificering. Nogle AI-systemer, der

udgør "uacceptable" risici, vil blive forbudt, mens en lang række AI-systemer med "høj risiko", der kan have en negativ indvirkning på menneskers sundhed, sikkerhed eller grundlæggende rettigheder, vil blive tilladt, men underlagt strenge krav for adgang til EU-markedet.

1. Regulering og relation til VoC

EU's tilgang til regulering af kunstig intelligens afspejler hovedsageligt en koordineret markedsøkonomi (CME). En CME prioriterer en stærk koordinering mellem stat, virksomheder og civilsamfund, som ses i reguleringens artikel 1, der fastslår formålet med forordningen at "etablere en ramme for kunstig intelligens og ændrer visse lovgivningsmæssige akter for at sikre, at AI udvikles, iagttages og anvendes på en sikker, ansvarlig og etisk måde i overensstemmelse med EU's værdier og grundlæggende rettigheder".

Tættere koordinering og stabilitet: Den tættere koordinering, der er kendetegnende for CME-modellen, kan bidrage til mere stabil vækst i EU. Som nævnt i artikel 22, der omhandler "Koordinering af tilsynet", er der etableret et europæisk AI-råd, der skal "bidrage til at fremme et ensartet anvendelsesniveau for denne forordning og til at sikre effektivt samarbejde mellem de nationale tilsynsmyndigheder". Dette kan skabe et mere forudsigeligt reguleringsmiljø for AI, hvilket kan føre til øget innovation, produktivitet og beskæftigelse.

Forebyggelse af arbejdsmarkedsudfordringer: Den koordinerede tilgang kan også bidrage til at forebygge negative effekter på arbejdsmarkedet. Artikel 28 om "Uddannelse og omstilling" adresserer dette ved at pålægge Kommissionen at "udvikle retningslinjer for medlemsstaterne om uddannelse og omstilling af arbejdstagere, der er berørt af AI-teknologiernes anvendelse". Dette kan bidrage til at sikre, at arbejdstagere har de nødvendige færdigheder til at navigere i en AI-drevet verden og dermed reducere risikoen for massearbejdsløshed og forringet trivsel på arbejdsmarkedet (European Commission, 2021).

2. Institutionelle Kontekster og AI-Regulering

De institutionelle kontekster i EU spiller en afgørende rolle i at forme tilgangen til AI-regulering. EU's AI Act afspejler dette ved at etablere et omfattende framework for involvering af forskellige aktører. Artikel 36 om "Samarbejde med NGO'er og interessenter" fastslår, at Kommissionen skal "oprette en ekspertgruppe for kunstig intelligens, der skal bestå af repræsentanter fra NGO'er, virksomheder, forskere og andre interessenter". Dette bidrager til en kollektiv og integreret tilgang,

der sikrer, at forskellige perspektiver tages i betragtning i udformningen af AI-politikken (European Commission, 2021).

3. Specifikke Reguleringsrammer for AI-Udvikling

EU's AI Act etablerer specifikke reguleringsrammer, der klassificerer AI-systemer baseret på risikoniveau og fastsætter overensstemmelseskrav for højrisiko AI-systemer. Artikel 9 om "Klassificering af AI-systemer" definerer fire kategorier af AI-systemer baseret på risikoen forbundet med deres anvendelse, hvor hver kategori har sine specifikke krav. For eksempel skal højrisiko AI-systemer overholde krav om gennemsigtighed, robusthed og menneskelig kontrol.

Balance mellem innovation og beskyttelse: EU's AI-lovgivning søger at skabe en balance mellem at fremme innovation og at beskytte borgernes rettigheder og sikkerhed. Dette afspejles i lovgivningens fokus på risikobaseret regulering, hvor AI-systemer klassificeres baseret på deres potentiale til at forårsage skade. Højrisiko AI-systemer er underlagt strengere krav, mens lavere risikofyldte systemer er underlagt mindre omfattende regulering (European Commission, 2021).

Forebyggende tilgang: EU's AI-lovgivning lægger vægt på en forebyggende tilgang til regulering. Dette indebærer at stille krav til AI-systemer på forhånd, inden de sættes i drift eller markedsføres. Dette skal sikre, at AI-systemer designes og udvikles på en måde, der minimerer risikoen for skade.

Forbud mod uacceptable praksisser: EU's AI-lovgivning forbyder visse AI-praksisser, der anses for at være uacceptable, f.eks. social scoring og brugen af AI til masseovervågning. Dette skal sikre, at AI-teknologier anvendes på en måde, der er i overensstemmelse med EU's værdier og grundlæggende principper.

Konformitetsvurdering: EU's AI-lovgivning kræver konformitetsvurdering af højrisiko AI-systemer, inden de kan sættes i drift eller markedsføres. Dette indebærer en vurdering af, om systemet opfylder alle de krav, der er fastsat i lovgivningen.

Artikel 5 i forordningen indeholder en liste over praksisser, der er forbudt i forbindelse med kunstig intelligens. Dette inkluderer anvendelse af AI til at manipulere folks adfærd på en måde, der kan skade dem eller underminere deres autonomi. Dette afspejler en koordineret tilgang, hvor der søges konsensus og samarbejde mellem medlemsstaterne.

Beskyttelse af forbrugere og arbejdstagere forskellige artikler i forordningen fokuserer på at beskytte forbrugere og arbejdstagere mod skadelige virkninger af AI-systemer. Dette er i tråd med en koordineret økonomis fokus på social beskyttelse og rettigheder.

Forskellige artikler af forordningen lægger vægt på international konvergens og bred accept samt internationale standarder. Dette indikerer et ønske om at koordinere AI-regulering på globalt plan og fremme EU's model som en standard for andre økonomier. Et forsøg på den såkaldte "Brussels Effect", for at påvirke landes lovgivning (Bradford, 2020).

USA

USA's executive order etablerer otte vejledende principper og prioriteter for udvikling og brug af AI. Disse principper omfatter sikkerhed, innovation, konkurrence, arbejdstøtte, hensyntagen til AI-bias og borgerrettigheder, forbrugerbeskyttelse og privatlivets fred. Denne ordre kommer i kølvandet på en stigende global bevidsthed om AI's indflydelse på samfundet og økonomien. AI's evner, der vokser eksponentielt, skaber en nødvendighed for at lede og forme teknologiens udvikling på en måde, der afspejler samfundets værdier og normer. Præsident Bidens ordre understreger behovet for en koordineret indsats på tværs af den føderale regering for at sikre, at AI-systemer er sikre og pålidelige, før de frigives til offentligheden. Dette inkluderer krav om, at udviklere af de mest kraftfulde AI-systemer deler deres data og anden kritisk information med den amerikanske regering.

1. Regulering og relation til VoC

USA's tilgang til AI-regulering kan ses som en blanding af LME- og CME-elementer. På den ene side er der en stærk tro på markedsmekanismer og innovationens kraft ved at opfordre til privat sektor-ledet AI-udvikling og begrænset direkte statslig intervention.

Reguleringen af AI afspejler dette ved at opstille rammer, der skal sikre, at AI-udviklingen sker på en måde, der er sikker for samfundet og økonomien, samtidig med at den opmuntrer til fortsat innovation og vækst. USA's tilgang til regulering af kunstig intelligens (AI) afspejler et liberalt markedsøkonomisk (LME) system, der prioriterer markedsfrihed og begrænset statslig intervention. I modsætning til EU's koordinerede markedsøkonomi (CME) model, lægger LME'er vægt på konkurrence, innovation og minimal regulering. Ordren prioriterer at fremme forskning og udvikling inden for AI for at sikre, at USA forbliver en global leder inden for denne teknologi. Ordren fokuserer på at fjerne unødvendige barrierer for AI-innovation og undgår at indføre strenge reguleringsrammer med et fokus på at sikre, at udvikling og anvendelse er i overensstemmelse

med den amerikanske forfatning. Den prioriterer privat sektors rolle i at drive AI-innovation, i overensstemmelse med den amerikanske tradition for et markedsdrevet system. Virksomheder opfordres til at udvikle og implementere AI-teknologier, mens regeringen spiller en støttende rolle. Der lægges vægt på selvregulering og samarbejde mellem virksomheder, forskere og regeringen. Dette afspejler en tro på, at markedet kan håndtere AI-udfordringer effektivt uden overdreven bureaukratisk indblanding. Dekretet anerkender betydningen af at beskytte intellektuel ejendomsret i AI-teknologier for at fremme innovation og investeringer. Dette er afgørende for at tiltrække talent og ressourcer til AI-forskning og -udvikling (The White House, 2023).

2. Institutionelle kontekster

Den institutionelle kontekst i USA, såsom en stærk privat sektor og en vægt på individuelle rettigheder, påvirker tilgangen til AI-regulering. Dette ses i dekretets fokus på at beskytte privatlivets fred, fremme lighed og borgerrettigheder. Dekretet vil understøtte forbrugere og arbejdere, samtidig med at innovation og konkurrence fremmes. Den amerikanske forfatning begrænser regeringens magt til at regulere virksomheder og beskytter individuelle friheder. Dette har haft en betydelig indflydelse på et mål om selvregulering og minimalt bureaukratisk tilsyn. Kongressen har beføjelse til at vedtage love, der regulerer AI, men har hidtil ikke gjort det i stor udstrækning. Kongressen kan dog føre tilsyn med regeringens aktiviteter inden for AI og bevilge midler til forskning og udvikling. De kan også indføre lovgivning, der adresserer specifikke bekymringer, f.eks. diskrimination eller misbrug af AI. Dette afspejler en vis tilbageholdenhed med at gribe ind i markedet og en tro på, at privat sektor bedst kan håndtere AI-udfordringer. De amerikanske domstole spiller en vigtig rolle i fortolkningen af love og reguleringer, der vedrører AI. Retssager kan bidrage til at afklare juridiske spørgsmål om AI, f.eks. ejerskab af data, privatlivets fred og intellektuel ejendomsret. Domstolenes afgørelser kan også have en betydelig indflydelse på udviklingen og brugen af AI-teknologier. Flere føderale agenturer har ansvar for at regulere specifikke aspekter af AI, f.eks. Federal Trade Commission (FTC) for forbrugerbeskyttelse og Food and Drug Administration (FDA) for medicinsk udstyr. Delstater og kommuner kan også vedtage deres egne love og reguleringer af AI. Dette kan føre til fragmentering af forskellige regler grundet de forskellige niveauer af governance, hvilket kan skabe udfordringer for virksomheder (The White House, 2023).

3. Specifikke reguleringsrammer

De specifikke reguleringsrammer for AI-udvikling i dekretet, modsat EU, kun de otte vejledende principper og prioriteter, herunder AI's sikkerhed og sikkerhed, gennemsigtighed, ansvarlighed og respekt for brugerrettigheder. Disse rammer er designet til at guide USA's agenturer og organisationer i deres AI-relaterede aktiviteter og beslutninger.

I erkendelse af de globale implikationer af AI opfordrer dekretet til at "etablere robuste internationale rammer for AI", hvilket antyder et skridt i retning af en koordineret tilgang. Dette stemmer overens med EU og taler for internationalt samarbejde for at tackle de udfordringer og muligheder, som AI giver (The White House, 2023).

Inklusivitet og mangfoldighed er kernen i dekretets politiske beslutningsproces. Direktivet understreger, at de udøvende afdelinger og agenturer (agenturer), når de udfører de handlinger, der er beskrevet i denne ordre, skal overholde disse principper, hvor det er relevant og i overensstemmelse med gældende lov, samtidig med at de tager hensyn til synspunkter fra andre agenturer, industrien, medlemmer af den akademiske verden, civilsamfundet, fagforeninger, internationale allierede og partnere og andre relevante organisationer, i det omfang det er muligt. Denne tilgang er i overensstemmelse med CME's karakteristik af at involvere en bred vifte af sociale partnere i beslutningstagningen, hvilket sikrer, at forskellige perspektiver former AI-politikker.

Dekretet siger udtrykkeligt, at "kunstig intelligens skal være sikker og tryk." For at nå dette mål er det nødvendigt at implementere robuste, pålidelige, gentagelige og standardiserede evalueringer af AI-systemer samt politikker, institutioner og, hvor det er relevant, andre mekanismer til at teste, forstå og mindske risici fra disse systemer, før de tages i brug. Dette direktiv kræver en omfattende og samarbejdsbaseret evalueringsproces, der afspejler de koordinerede markedsøkonomiers (CME'er) præference for ikke-markedsmæssige interaktioner for at opnå koordinering og styre risici. Dekretets fokus på at udvikle "standarder, værktøjer og tests til AI-systemer" signalerer en forpligtelse til langsigtede investeringer i AI's sikkerhed og sikkerhedsinfrastruktur. Denne tilgang afspejler den tålmodige kapitaltilgang, der er typisk for CME'er, hvor investeringer ikke udelukkende er drevet af øjeblikkelige økonomiske afkast, men af bæredygtig vækst og innovation (The White House, 2023).

Kina

1. Regulering og relation til VoC

Kinas tilgang til regulering af kunstig intelligens (AI) kan ses som en afspejling af landets statskapitalistiske system, hvor staten spiller en central rolle i økonomisk styring og regulering, samtidig med at den tillader og fremmer privat virksomhed inden for visse rammer. Den statskapitalistiske model tillader Kina at udøve stram kontrol med strategiske sektorer, herunder teknologi og AI, hvilket er tydeligt i landets lovgivningsmæssige tiltag. Kina har udviklet en omfattende ramme for AI-styring, som inkluderer strategier, initiativer og nøgleovervejelser, der passer ind i landets større databeskyttelsesramme. Dette understøtter ideen om, at Kina søger at balancere innovation og kontrol, hvilket er karakteristisk for en statskapitalistisk tilgang.

Kinas regulering af AI er også præget af en pro-vækst holdning, som det ses i de midlertidige foranstaltninger og en banebrydende domstolsafgørelse, der giver ophavsretlig beskyttelse til AI-output. Dette viser en vilje til at fremme teknologisk udvikling, mens man sikrer, at staten bevarer den nødvendige kontrol for at beskytte nationale interesser og sikkerhed. Samtidig med at Kina anerkender risikoen ved at gå glip af mulighederne inden for AI-teknologi, er der et klart fokus på at etablere de nødvendige betingelser for at indhente teknologifronten, herunder at udnytte generativ AIs potentiale til at øge økonomisk produktivitet og fremme videnskabelige fremskridt i andre discipliner (DigiChina, 2023)

Den statslige indflydelse er yderligere tydelig i forslaget fra den statsaffilierte Chinese Academy of Social Sciences (CASS), som foreslår oprettelsen af en statslig myndighed, der skal regulere AI, og implementeringen af en negativliste-mekanisme, der ville underlægge højrisiko-AI-systemer strengere regulering. Dette afspejler en tilgang, hvor staten ikke kun regulerer AI gennem eksisterende love og politikker, men også aktivt former fremtidens AI-landskab gennem nye lovgivningsinitiativer (DigiChina, 2023).

Samlet set afspejler Kinas AI-regulering landets statskapitalistiske system ved at kombinere en stærk statslig styring med en fremme af privat innovation inden for AI. Dette skaber en unik dynamik, hvor staten både er en facilitator og en regulator, hvilket er i tråd med den overordnede

økonomiske og politiske strategi i Kina. Ved at regulere AI på denne måde søger Kina at sikre, at teknologien udvikler sig på en måde, der understøtter landets økonomiske mål og politiske stabilitet, samtidig med at det giver plads til innovation og vækst inden for den private sektor.

2. Institutionelle kontekster

Kinas institutionelle kontekster som en statskapitalistisk økonomi har en dybtgående indflydelse på tilgangen til AI-regulering. I et system, hvor staten spiller en dominerende rolle som ejer eller investor-aktionær, er der en unik fusion mellem regeringens politiske mål og økonomiske aktiviteter. Dette system kan både fremme og hæmme udviklingen af AI-teknologier og deres regulering. På den ene side kan statskapitalismen muliggøre betydelige investeringer i AI-forskning og -udvikling, hvilket kan accelerere teknologisk innovation og implementering. Den kinesiske regerings evne til at koordinere og dirigere ressourcer mod strategiske sektorer kan også facilitere en mere omfattende og ensartet regulering af AI.

Den tætte forbindelse, i form af den høje grad af korporatisme mellem stat og økonomi, kan føre til en regulering, der er mere fokuseret på at opretholde statens kontrol og mindre på at fremme konkurrence og innovation. Reguleringen kan blive brugt som et redskab til at styrke indenlandske virksomheder på bekostning af udenlandske konkurrenter, hvilket kan begrænse internationalt samarbejde og udveksling af viden. Desuden kan den statslige indblanding i økonomien og virksomhedernes drift skabe uklarhed omkring ejerskab og ansvar, hvilket kan gøre det vanskeligt at fastlægge ansvarlighed i tilfælde af misbrug af AI eller fejl (DigiChina, 2023).

Kinas AI-lovgivning afspejler en ambition om at blive en global leder inden for AI. Lovgivningen lægger vægt på at fremme innovation samtidig med, at den sikrer national sikkerhed og social stabilitet. Dette kan ses som et forsøg på at balancere behovet for teknologisk fremskridt med regeringens ønske om at bevare kontrol. Reguleringen søger at etablere et rammeværk for etisk AI-udvikling, der respekterer brugernes rettigheder og fremmer offentlig tillid til teknologien.

Kinas tilgang til AI-regulering er et resultat af en kompleks interaktion mellem statens økonomiske interesser, politiske mål og det internationale miljø. Mens statskapitalismen giver regeringen værktøjer til at forme AI-udviklingen, rejser det også spørgsmål om, hvordan disse værktøjer anvendes, og hvilken indvirkning de vil have på både det kinesiske samfund og det globale AI-landskab. Det er afgørende for internationale observatører og aktører inden for AI at forstå disse

dynamikker for at navigere i og samarbejde med Kina på dette vigtige og hurtigt udviklende område.

3. Specifikke reguleringsrammer

Kinas reguleringsrammer for AI-udvikling er et komplekst system, der afspejler landets ambitioner om at være en global leder inden for kunstig intelligens.

Kinas tilgang til AI-regulering adresserer specifikke udfordringer såsom AI-drevne anbefalingssystemer, dybsynteseværktøjer og brug af ansigtsgenkendelse af ikke-statslige agenturer.

Ifølge Artikel 1 er lovens formål at regulere forskning og udvikling, levering og brug af kunstig intelligens (AI), for at beskytte national suverænitet, udvikling og sikkerhed, fremme sikker udvikling af AI og beskytte legitime rettigheder og interesser for individer og organisationer. Loven specificerer i Artikel 2, at den gælder for forskning og udvikling, levering og brug af AI samt regulering af AI inden for Folkerepublikken Kinas grænser. Aktiviteter relateret til AI, der udføres uden for landets territorium, men som påvirker eller kan påvirke national sikkerhed, offentlige interesser eller legitime rettigheder og interesser for individer eller organisationer i Kina, er også underlagt denne lov (DigiChina, 2023).

Artikel 3 fremhæver, at staten skal koordinere udvikling og sikkerhed, holde fast i kombinationen af at fremme innovation og styring i henhold til loven og implementere inklusiv og forsigtig regulering. Artikel 4 understreger et menneskecentreret princip, hvor aktiviteter relateret til AI skal være menneskecentrerede og dirigere intelligens til det gode, sikre kontinuerlig menneskelig overvågning og kontrol med AI, med det ultimative mål altid at fremme menneskehedens velfærd. Sikkerhedsprincipperne i Artikel 5 kræver, at de, der forsker og udvikler, leverer eller bruger AI, skal træffe nødvendige foranstaltninger for at sikre sikkerheden af forskning og udvikling, levering og brug af AI og de relaterede netværksdata (DigiChina, 2023).

Artikel 6 kræver åbenhed, gennemsigtighed og forklarbarhed fra dem, der forsker og udvikler, leverer eller bruger AI, og at de skal vedtage nødvendige politikker for at sikre, at formålet, principperne og effekterne af AI-forskning og -udvikling samt brug af AI overholder ovenstående. Endelig fastsætter Artikel 7 et ansvarlighedsprincip, hvor de, der er engageret i forskning og udvikling, levering og brug af AI, skal være individuelt ansvarlige for deres respektive aktiviteter.

Disse artikler udgør kernen i den foreslåede regulering, som søger at balancere innovation og sikkerhed, samtidig med at den beskytter borgernes rettigheder og fremmer en etisk tilgang til AI-udvikling. Det er tydeligt, at Kina tager et omhyggeligt skridt mod at forme en fremtid, hvor AI er en integreret og velreguleret del af samfundet, med en lovgivning, der afspejler landets prioriteter og værdier i den teknologiske æra. Den kinesiske regulering er et eksempel på, hvordan lovgivere kan tilpasse sig den hurtige udvikling inden for AI og sikre, at teknologien udvikles ansvarligt og med menneskelige interesser i centrum. Det er en indikation på, at Kina ønsker at være i forkant med at definere juridiske rammer for AI, hvilket kan have vidtrækkende konsekvenser for global AI-styring (DigiChina, 2023).

Det er tydeligt, at Kina tager AI-regulering alvorligt og har implementeret en omfattende regulering, der spænder over flere jurisdiktioner og reguleringsniveauer. Dette system er designet til at sikre, at AI-udviklingen sker på en måde, der er i overensstemmelse med national sikkerhed, sociale offentlige interesser og beskyttelse af borgernes rettigheder. Disse regler er specifikke for Kina, kan deres indflydelse give genlyd globalt og inspirere til lignende strategier i andre lande, især lande med lignende komparative markedsøkonomiske elementer.

Opsamling på analyse

I liberale markedsøkonomier (LME) som USA, er der en tendens til at favorisere markedsmekanismer og begrænse statslig intervention, hvilket afspejles i den amerikanske regerings Executive Order on AI, der fremhæver innovation og konkurrence. I modsætning hertil står EU, som repræsenterer en mere koordineret markedsøkonomi (CME), hvor der er en større balance mellem markeds kræfter og regulering på EU-niveau, som det ses i AI Act. Dette lovgivningsværk søger at sikre fair konkurrence og forbrugerbeskyttelse, samtidig med at det fremmer etisk og socialt ansvar i AI-udviklingen.

Kina adskiller sig ved at fremme privat sektor-ledet innovation inden for en statskapitalistisk model, hvor staten spiller en central rolle i at styre og fremme nationalt teknologisk lederskab, som det er angivet i deres National AI Development Plan. Mens alle tre økonomier anerkender vigtigheden af markedsmekanismer for at drive innovation, varierer graden af statslig involvering og regulering betydeligt.

EU's tilgang til AI-regulering inkluderer koordinering og samarbejde på både nationalt og EU-niveau, hvilket er afgørende for at udvikle sammenhængende strategier og reguleringer. AI Act fremhæver behovet for detaljerede og præventive reguleringsrammer, herunder klassificering af AI-systemer baseret på risiko, og krav til risikovurdering og -mitigering. Der lægges også vægt på etisk ansvar og socialt ansvar, med etiske retningslinjer for AI-udvikling og -brug, samt krav om gennemsigtighed og redegørelse for AI-systemer.

I kontrast hertil har USA og Kina et mere begrænset fokus på CME-elementer. USA fokuserer på markedstilpasning og efterfølgende regulering, mens Kina prioriterer centraliseret styring. Disse forskelle i tilgang til AI-regulering afspejler dybere ideologiske og strukturelle forskelle mellem LME og CME markedsøkonomier, som har betydelige implikationer for global konkurrence og innovation inden for AI. Det er anskueligt, at mens markedsmekanismer er centrale for AI-udviklingen, er der et stigende behov for at overveje de etiske og sociale konsekvenser af AI, hvilket EU's reguleringstilgang søger at adressere i højere grad.

LME-elementer

Både USA, EU og Kina fremhæver privat sektor-ledet innovation i AI-udviklingen, men med forskellige tilgange:

I USA er LME-elementerne karakteriseret ved en stærk tro på markedets evne til selvregulering og innovation uden tung statslig indgriben. Dette kommer til udtryk gennem politikker som Executive Order on AI, hvor Sektion 5 fremhæver en præference for markedsmekanismer frem for statslig intervention. Den amerikanske tilgang til AI er rodfæstet i en laissez-faire økonomisk filosofi, hvor private virksomheder forventes at lede an i teknologisk udvikling, og staten primært fungerer som en facilitator snarere end en regulator.

I EU's tilfælde afspejler LME-elementerne en mere nuanceret synsvinkel, hvor der er en større vægt på regulering for at sikre fair konkurrence og forbrugerbeskyttelse. Dette ses i AI Act, hvor Artikel 3 lægger op til en balanceret tilgang, der tillader markedsmekanismer at operere inden for et reguleret miljø. EU's model søger at harmonisere det indre marked, samtidig med at man sikrer, at AI-teknologi udvikles og anvendes på en måde, der er i overensstemmelse med europæiske værdier og standarder.

Kinas LME-elementer er indlejret i en statskapitalistisk kontekst, hvor staten spiller en mere aktiv rolle i økonomien. Kina støtter privat sektor-ledet innovation, men inden for en ramme, der er stærkt påvirket af statens strategiske mål. Dette indebærer en betydelig grad af statslig styring og støtte, hvilket afspejler en model, hvor markedsmekanismerne er til stede, men underlagt en stærkere statslig kontrol og retning.

CME-elementer

I USA er CME-elementer mindre fremtrædende, da landet traditionelt favoriserer en liberal markedsökonomi, hvor markeds kræfterne har større frihed til at forme økonomien. Dette betyder, at selvom der er visse former for koordinering, såsom standardisering og patentering, som fremmer samarbejde mellem virksomheder, er der en tendens til at regulere markedet med en mere reaktiv tilgang. Reguleringen sker ofte som respons på udviklinger i stedet for at være proaktiv, hvilket kan ses i den seneste tids håndtering af tech-giganterne og anvendelse af konkurrencelovgivning. EU står i kontrast til både USA og Kina med sin unikke tilgang til CME. EU's økonomiske model er karakteriseret ved en højere grad af koordinering og regulering, som søger at balancere markeds kræfterne med sociale og etiske overvejelser. EU's AI Act, fremhæver vigtigheden af koordinering og samarbejde mellem staten, virksomheder og civilsamfundet for at sikre, at AI-udviklingen sker på en ansvarlig og etisk måde. EU's model fremmer også socialt ansvar og etiske standarder, som er afgørende for at sikre, at teknologien tjener samfundets bedste interesser. (AI Act, artikel 5, 13)

I Kina, derimod, er CME-elementerne indlejret i en stærkere statsstyret økonomisk model. Staten spiller en central rolle i koordineringen af økonomiske aktiviteter, hvilket inkluderer direkte indgriben i markederne og omfattende planlægning. Dette ses tydeligt i Kinas tilgang til AI-udvikling, hvor regeringen har sat ambitiøse mål og skabt omfattende nationale planer for at blive en førende magt inden for AI. Denne centraliserede tilgang sikrer en høj grad af koordinering og samarbejde på tværs af forskellige sektorer, men rejser også spørgsmål om innovationens frihed og individuelle virksomheders autonomi. Alle AI modeller skal godkendes på central hånd af Kinas etpartisystem styret af CCP, om de enkelte modeller overholder regimets marxistiske ideologi og den såkaldte Xi Jinping tankegang.

USA	Kina	EU
Fokus på innovation og minimal statslig indgriben.	Fokus på national kontrol og teknologisk suverænitet.	Fokus på konsensus og harmonisering af standarder.
Tillid til selvregulering: Mere frivillige standarder og private initiativer.	Topstyret regulering: Strengte databegrænsninger og fokus på national sikkerhed.	Forsigtig tilgang: Prioritering af databeskyttelse og etiske principper.
Markedsbaseret tilgang med begrænset statslig intervention.	Statsdrevet tilgang med fokus på national sikkerhed og social stabilitet.	Koordineret tilgang med inddragelse af forskellige interesser.
Beføjelser til agenturer til at beskytte forbrugere og arbejdstagere.	Regimet vurderer og godkender AI-værktøj ud fra om de overholder partiets værdier.	Klassificering af AI-systemer baseret på risiko, krav til højrisiko-systemer og forbud mod visse anvendelser.
Stærk privat sektor og individuelle rettigheder.	Staten spiller en central rolle i udvikling og regulering.	Koordinering og samarbejde nationalt og på EU-niveau.
Bekymring for databeskyttelse, bias og diskrimination.	Kontrol med strategiske sektorer, herunder teknologi og AI.	Balance mellem innovation og beskyttelse af borgernes rettigheder.

Diskussion

I lyset af specialets analyse og resultater, kan det konkluderes, at VoC-teorien har været et værdifuldt redskab til at forstå og forklare de forskellige tilgange til regulering af AI i EU, USA og Kina. Diskussionen vil bearbejde analysen og sammenholde resultaterne heraf med specialets 4 hypoteser baseret ud fra VoC-tilgangen.

Hypotese 1

Specialets resultater understøtter i høj grad Hypotese 1.

Analysen af EU's AI Act bekræfter, at reguleringen afspejler en koordineret markedsøkonomi (CME), hvor der lægges vægt på sociale hensyn og forsigtighedsprincippet. Dette ses blandt andet i forbuddet mod visse anvendelser af AI, der anses for at være højrisiko eller etisk problematiske, samt i kravene om gennemsigtighed og menneskelig overvågning af højrisikosystemer. Disse bestemmelser afspejler en beskyttende tilgang til regulering, der er karakteristisk for CME'er. Samtidig viser AI Act's fleksible elementer, såsom regulatoriske sandkasser og undtagelser for forskning og innovation, en vis grad af tilpasningsevne og en erkendelse af behovet for at fremme teknologiudvikling. Dette understøtter antagelsen om, at EU's tilgang til AI-regulering er præget af en vis grad af fleksibilitet og kompromis mellem forskellige interesser.

Analysen af USA's dekret om AI bekræfter også hypotesen. Dekretet lægger vægt på at fremme innovation og minimere statslig indblanding, hvilket er karakteristisk for en liberal markedsøkonomi (LME). Den understreger principper som ansvarlig innovation, konkurrence og samarbejde, og opfordrer til, at agenturer undgår unødige begrænsninger på AI-udviklingen. Dog åbner dekretet også op for en mere proaktiv tilgang, hvis der opstår alvorlige sikkerheds- eller etiske problemer i forbindelse med AI-anvendelser. Dette indikerer, at selvom USA's tilgang er baseret på LME-principper, er der en potentiel åbning for øget statslig involvering i fremtiden.

Analysen af Kinas AI-regulering bekræfter hypotesen om en statskapitalistisk model. Reguleringen fokuserer på national sikkerhed og social stabilitet, hvilket afspejler Kinas politiske system og værdier. Staten spiller en central rolle i at definere og styre udviklingen og anvendelsen af AI. Dog viser analysen også, at Kina anerkender vigtigheden af markedsbaseret innovation og internationalt samarbejde for at fremme økonomisk vækst. Dette ses i bestemmelserne om at støtte forskning og udvikling inden for AI og fremme samarbejde med internationale partnere. Kinas tilgang til AI-regulering er således ikke udelukkende statskapitalistisk, men inddrager også elementer af markedsbaseret innovation og internationalt samarbejde.

Samlet set understøtter specialets resultater Hypotese 1 og viser, at AI-reguleringen i EU, USA og Kina afspejler deres respektive kapitalismemodeller. Dog er der også nuancer i de enkelte landes tilgange, hvilket indikerer, at reguleringen ikke blot er et resultat af den overordnede kapitalismemodel, men også påvirkes af andre faktorer som historiske begivenheder, politiske diskurser og specifikke nationale interesser.

Hypotese 2

Specialets analyse giver delvis støtte til Hypotese 2.

I EU understøttes hypotesen delvist. EU's institutionelle rammer, karakteriseret ved en stærk tilstedeværelse af civilsamfundsaktører og en tradition for samarbejde mellem stater, har resulteret i en relativt omfattende og inkluderende AI-regulering i form af AI Act. Forslaget inddrager en bred vifte af interessenter og lægger vægt på grundlæggende rettigheder og etiske principper, hvilket stemmer overens med forventningerne til en koordineret markedsøkonomi (CME). Dog har den lange og komplekse lovgivningsproces i EU, med forhandlinger og kompromisser mellem medlemslandene, ført til en vis udvanding af nogle af de oprindelige forslag i AI Act. Dette viser,

at selvom EU's institutionelle rammer muliggør en inkluderende reguleringsproces, kan den også føre til forsinkelser og udvanding af bestemmelserne.

I USA understøttes hypotesen også delvist. Landets institutionelle rammer, med en stærk vægt på individuelle rettigheder og markedsfrihed, har resulteret i en mere begrænset statslig indblanding i AI-reguleringen, primært i form af den præsidentielle bekendtgørelse. Dekretet fokuserer på at fremme innovation og konkurrence, men anerkender også behovet for at beskytte visse samfundsmæssige interesser. Dog er den politiske proces omkring AI-regulering i USA præget af lobbyisme fra teknologivirksomheder og interessekonflikter mellem forskellige aktører, hvilket kan begrænse reguleringens omfang og effektivitet. Dette viser, at selvom de institutionelle rammer i en LME begrænser statslig indblanding, kan de også føre til en situation, hvor stærke private interesser får stor indflydelse på reguleringen.

I Kina bekræftes hypotesen kun delvist. Kinas centraliserede statsmagt og partidominerede system har ført til en hurtig og effektiv implementering af AI-regulering. Dog viser analysen af Kinas 24 retningslinjer for kunstig intelligens, at der er en vis grad af inddragelse af interessenter, herunder eksperter og virksomheder, i udformningen af reguleringen. Dette tyder på, at selvom staten har en stærk rolle i Kinas AI-regulering, er der også en vis åbenhed over for input fra andre aktører. Dog er der fortsat bekymringer om den begrænsede inddragelse af civilsamfundet og uafhængige eksperter, hvilket kan øge risikoen for utilsigtede konsekvenser og begrænse teknologiens potentiale.

Samlet set understøtter specialets resultater delvist Hypotese 2, men viser også, at virkeligheden er mere kompleks end de idealtypiske kategorier i VoC-teorien. Både EU og USA viser tegn på en vis grad af fleksibilitet og tilpasningsevne i deres AI-regulering, mens Kina, på trods af sin statsdominerede model, også viser tegn på inddragelse af interessenter. Dette peger på behovet for at nuancere VoC-tilgangen og inddrage andre faktorer, såsom historiske begivenheder, politiske diskurser og aktørernes strategier, for at opnå en mere fyldestgørende forståelse af, hvordan AI-reguleringen formes og implementeres i forskellige lande.

Hypotese 3

Specialets analyse giver delvis støtte til Hypotese 3.

I EU understøttes hypotesen, idet AI Act foreslår en risikobaseret tilgang til AI-regulering, hvor højrisiko-AI-systemer underkastes strenge krav, herunder krav til risikostyring, datakvalitet og menneskelig overvågning. Dog indeholder AI Act også elementer af fleksibilitet, der tager højde for behovet for innovation og teknologiudvikling. For eksempel tillader forordningen "regulatoriske sandkasser", hvor virksomheder kan teste nye AI-systemer i et kontrolleret miljø, og der er mulighed for undtagelser for forskning og innovation. Disse elementer afspejler en balance mellem en præventiv tilgang og behovet for at fremme innovation inden for AI.

I USA understøttes hypotesen kun delvist. Dekretet om AI lægger vægt på frivillige retningslinjer og standarder, hvilket indikerer en mere reaktiv tilgang til regulering. Dog åbner dekretet også op for muligheden for en mere proaktiv tilgang, hvis der opstår alvorlige sikkerheds- eller etiske problemer i forbindelse med AI-anvendelser. Dette tyder på, at USA's tilgang ikke udelukkende er reaktiv, men at der er en vis åbenhed over for at stramme reguleringen, hvis det bliver nødvendigt.

Kinas tilgang til AI-regulering er, som forventet, en blanding af statslig kontrol og støtte til udvalgte industrier. Loven forbyder visse anvendelser af AI, der kan true national sikkerhed eller social stabilitet, men tilskynder også til udvikling af AI-systemer, der fremmer økonomisk vækst og social velfærd. Der er dog også tegn på en vis grad af fleksibilitet og eksperimentering i Kinas tilgang, især inden for særlige økonomiske zoner, hvor virksomheder får større frihed til at udvikle og teste nye teknologier. Dette viser, at selvom staten spiller en central rolle i Kinas AI-regulering, er der også en erkendelse af behovet for at fremme innovation og økonomisk vækst.

Samlet set bekræfter analysen delvist Hypotese 3, men viser også, at virkeligheden er mere kompleks end de idealtypiske kategorier i VoC-teorien. Alle tre viser tegn på en vis grad af fleksibilitet og tilpasningsevne i deres AI-regulering, selvom de grundlæggende tilgange er forskellige.

Hypotese 4

Analyse af EU's AI Act understøtter i høj grad Hypotese 4. EU's tilgang til AI-regulering er præget af en risikobaseret tilgang med fokus på etik, sikkerhed og gennemsigtighed, hvilket stemmer overens med de værdier, der forventes i en koordineret markedsøkonomi (CME). Eksempelvis forbyder AI Act visse anvendelser af AI, der anses for at være højrisiko eller etisk problematiske, og stiller strenge krav til gennemsigtighed og menneskelig overvågning af højrisikosystemer. Dette kan ses som et forsøg på at styre innovationen i en retning, der er i overensstemmelse med europæiske værdier og standarder, og som prioriterer beskyttelse af borgernes rettigheder og sikkerhed. Dog kan denne tilgang potentielt hæmme innovation på kort sigt, da den stiller strenge krav til visse AI-systemer. På længere sigt kan den risikobaserede tilgang dog fremme tillid til teknologien og skabe et mere stabilt og forudsigeligt reguleringsmiljø, hvilket kan være en fordel for europæiske virksomheder.

Analysen af USA's bekendtgørelse om AI understøtter delvist Hypotese 4. Dekretet lægger vægt på at fremme innovation og fjerne barrierer for AI-udvikling, hvilket stemmer overens med forventningerne til en liberal markedsøkonomi (LME). Dog er der også elementer i dekretet, der peger i retning af en mere proaktiv tilgang til regulering, især når det gælder beskyttelse af privatlivets fred, borgerrettigheder og forbrugerbeskyttelse. Dette indikerer, at selvom USA's tilgang primært er markedsdrevet, er der en erkendelse af behovet for at adressere de potentielle risici og negative konsekvenser af AI. Spørgsmålet er, om den nuværende tilgang er tilstrækkelig til at håndtere disse risici på længere sigt, eller om der vil være behov for en mere proaktiv og risikobaseret regulering, som vi ser det i EU.

Specialets resultater bekræfter i høj grad Hypotese 4 i forhold til Kina. Kinas statslige styring og fokus på national sikkerhed og social stabilitet afspejles i landets AI-regulering. Dette giver staten mulighed for at målrette investeringer og ressourcer mod strategisk vigtige AI-områder, hvilket kan fremme innovation inden for disse specifikke områder. Dog kan den statslige kontrol og overvågning også begrænse den overordnede innovation og skabe etiske dilemmaer i forhold til overvågning og kontrol. Derudover kan den manglende inddragelse af civilsamfundet og uafhængige eksperter i beslutningsprocessen føre til en ulige fordeling af fordele og ulemper ved AI-udviklingen.

Samlet set understøtter specialets resultater Hypotese 4, men med visse nuancer: EU's tilgang til AI-regulering primært fokuserer på at styre innovationen i en etisk og ansvarlig retning, så er USA's tilgang mere fokuseret på at fremme innovation og økonomisk vækst, men med en erkendelse af behovet for at beskytte visse samfundsmæssige interesser. Kinas tilgang er præget af en stærk statslig styring og kontrol, hvilket kan have både positive og negative konsekvenser for innovation og udvikling af AI.

Opsamling

VoC-teoriens begreber og ramme kan anvendes til at analysere AI-regulering i EU, USA og Kina, men den har også sine begrænsninger, især når det gælder forståelsen af AI, der er en ny og hastigt udviklende teknologi. VoC-tilgangen (Hall & Soskice, 2001) er primært udviklet til at analysere etablerede industrier og økonomiske sektorer, og den tager ikke fuldt ud højde for de unikke udfordringer og muligheder, som AI bringer med sig. For det andet fokuserer VoC-tilgangen primært på nationale institutionelle rammer og overser den stigende betydning af globale aktører og processer i udformningen af AI-regulering (Palley et al., 2023). Tech-giganter som Google, Microsoft, Amazon og Meta har en enorm indflydelse på AI-landskabet på grund af deres store ressourcer, dataadgang og markedsdominans. Deres lobbyvirksomhed og strategiske partnerskaber med regeringer kan påvirke reguleringen i en retning, der gavner deres egne interesser, hvilket kan være i modstrid med de bredere samfundsmæssige interesser. Endelig tager VoC-tilgangen ikke højde for den hurtige teknologiske udvikling inden for AI og de etiske og sociale udfordringer, som dette medfører. AI-teknologien udvikler sig hurtigere, end lovgivere og regulatorer kan følge med, hvilket skaber en skævvridning i regulering på forskellige niveauer af governance. Derudover rejser AI en række etiske spørgsmål, såsom bias, diskrimination og ansvarlighed, som ikke altid kan løses inden for de eksisterende institutionelle rammer. På trods af disse begrænsninger har VoC-tilgangen vist sig at være et værdifuldt redskab til at analysere de forskellige tilgange til AI-regulering i USA, Kina og EU.

Normativt set har specialets resultater flere implikationer for, hvordan man bør regulere AI fremadrettet. Specialet viser, at der ikke findes en universel model for AI-regulering. Reguleringen bør i stedet tilpasses de specifikke institutionelle rammer og værdier i forskellige lande. En risikobaseret tilgang som EU's kan være passende i koordinerede markedsøkonomier (CME), mens en mere fleksibel tilgang kan være bedre i liberale markedsøkonomier (LME). Internationalt

samarbejde i forhold til regulering kræver derfor en høj grad af landenes villighed til at bøje sig mod hinanden, især fra LME'er som USA, da LME'er i mindre grad vil regulere kunstig intelligens end CME'er og statskapitalistiske modeller.

Specialet understreger vigtigheden af at finde en balance mellem at fremme innovation og beskytte samfundsmæssige interesser. En for streng regulering kan hæmme innovation, mens en for lempelig regulering kan føre til uønskede konsekvenser. Reguleringen skal være dynamisk og tilpasningsdygtig for at følge med den hurtige teknologiske udvikling inden for AI. Løbende evaluering og revision er nødvendigt. Specialet kunne styrkes ved at inddrage flere teoretiske perspektiver for at opnå en mere nuanceret og helhedsorienteret forståelse af AI-reguleringen, samt en diskussion af, hvordan de tre cases forskellige tilgange til AI-regulering kan påvirke det internationale samarbejde og konkurrencen inden for AI. Det kunne også være relevant at undersøge, hvordan AI-reguleringen i de tre cases påvirker og påvirkes af andre globale aktører, såsom internationale organisationer og tech-giganter.

Konklusion

Specialet har undersøgt, hvordan reguleringen af kunstig intelligens (AI) i USA, Kina og EU kan forstås i et varieties of capitalism (VoC) perspektiv. Gennem anvendelse af VoC-teori i en komparativ analyse af policy-dokumenter og med inddragelse af historisk institutionalisme er det blevet påvist, at der er en klar sammenhæng mellem landenes institutionelle rammer og deres tilgang til AI-regulering. EU's risikobaserede tilgang, med fokus på etik, sikkerhed og gennemsigtighed, afspejler en koordineret markedsøkonomi (CME), hvor sociale hensyn prioriteres højt. USA's mere fleksible og markedsdrevne tilgang, der lægger vægt på innovation og økonomisk vækst, er karakteristisk for en liberal markedsøkonomi (LME). Kinas statsdominerede model, med fokus på national sikkerhed og social stabilitet, afspejler landets unikke politiske og økonomiske system. Studiet bekræfter således VoC-teoriens grundlæggende antagelse om, at forskellige kapitalismemodeller fører til forskellige reguleringsstrategier. Dog viser analysen også, at der er nuancer og kompleksiteter, som VoC-teorien ikke fuldt ud indfanger: Kinas tilgang til AI-regulering ikke udelukkende statskapitalistisk, men inddrager også elementer af markedsbaseret innovation. Ligeledes er EU's risikobaserede tilgang ikke så stringent, som VoC-teorien måske kunne forudsige, da den også indeholder fleksible elementer for at fremme

innovation.

Specialet har også vist, at historiske begivenheder og politiske diskurser spiller en væsentlig rolle i udformningen af AI-regulering. EU's historiske fokus for databeskyttelse og fokus på borgerrettigheder har påvirket udformningen af EU's regulering af kunstig intelligens, hvorimod USA's historiske vægt på markedsfrihed og innovation har formet landets mere liberale tilgang. I Kina har det historiske fokus på national sikkerhed og social stabilitet spillet en afgørende rolle i AI-strategi. AI-regulering ikke blot er et spørgsmål om tekniske og juridiske overvejelser, men i særdeleshed et politisk spørgsmål med vidtrækkende konsekvenser for verdenssamfundet.

Specialet bidrager til en voksende litteratur om AI-regulering og VoC. Ved at anvende VoC-teorien og historisk institutionalisme på et nyt og relevant område, giver specialet en dybere forståelse af, hvordan forskellige kapitalistiske modeller påvirker reguleringen af denne revolutionære teknologi. Resultaterne har implikationer for både den akademiske debat og den politiske praksis, da de viser, at der ikke er én universel model for AI-regulering, men at reguleringen skal tage højde for den specifikke kontekst og de værdier, der kendetegner forskellige kapitalistiske modeller. VoC-teorien, selvom den er et nyttigt redskab, kan ikke forklare alle aspekter af AI-reguleringen. Den hurtige teknologiske udvikling, globaliseringens indvirkning og de etiske og sociale udfordringer ved AI kræver, at der inddrages andre teoretiske perspektiver og empiriske undersøgelser for at opnå en mere nuanceret og helhedsorienteret forståelse af, hvordan AI-reguleringen formes og implementeres. Yderligere forskning er nødvendig for at undersøge de langsigtede konsekvenser af forskellige AI-reguleringsmodeller og for at udvikle mere robuste og adaptive reguleringsrammer, der kan imødegå de unikke udfordringer og muligheder, som AI-teknologien bringer med sig. Specialets fokus på policy-dokumenter som primær datakilde har begrænsninger, da det ikke giver indsigt i den faktiske implementering og praksis af AI-regulering. Fremtidig forskning bør derfor supplere dokumentanalysen med andre metoder, såsom interviews med relevante aktører og casestudier af specifikke AI-anvendelser, for at opnå en mere dybdegående forståelse af, hvordan AI-regulering fungerer i praksis.

Bibliografi

- Amable, B. (2003, December 4). The Diversity of Modern Capitalism. OUP Academic. <https://doi.org/10.1093/019926113X.001.0001>
- Associated Press. (2021, May 1). Putin: Leader in artificial intelligence will rule world. Retrieved April 3, 2024, <https://apnews.com/article/bb5628f2a7424a10b3e38b07f4eb90d4>
- Bradford, A. (2020) The Brussels Effect: How the European Union Rules the World <https://doi.org/10.1093/oso/9780190088583.001.0001>
- Bulmer, S. (2023). Institutionalism. The Elgar Companion to the European Union, 27–36. <https://doi.org/10.4337/9781800883437.00010>
- De Vaus, D. A. (2001). Research design in social research. SAGE.
- DigiChina (2023). Translation: Artificial Intelligence Law, Model Law v. 1.0. <https://digichina.stanford.edu/work/translation-artificial-intelligence-law-model-law-v-1-0-expert-suggestion-draft-aug-2023/>
- DigiChina (2017). Full Translation: China's 'New Generation Artificial Intelligence Development Plan' (2017). <https://digichina.stanford.edu/work/full-translation-chinas-new-generation-artificial-intelligence-development-plan-2017/>
- European Commission. (2021). Fostering a European approach to artificial intelligence. Office for Official Publications of the European Communities.
- European Commission. (2023). Proposal for a regulation of the European Parliament and the Council laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act) and amending certain Union legislative acts.
- European Parliament, European Council and European commission. (2022). European declaration on digital rights and principles. Office for Official Publications of the European Communities.
- Hall, P. A., & Soskice, D. W. (Eds.). (2001). *An Introduction to Varieties of Capitalism: The Institutional Foundations of Comparative Advantage*. Oxford University Press.
- Gorski, P. S. (2013). "What is Critical Realism? And Why Should You Care?" Contemporary Sociology: A Journal of Reviews, 42(5), 658–670. <https://doi.org/10.1177/0094306113499533>
- Klitgaard, M. B., & Roederer-Rynning, C. (2009). Kapitalismens variationer og EU's østudvidelse. *Politica - Tidsskrift for Politisk Videnskab*, 41(1), 46-67.

- Maindonald, J. H. (2011). Qualitative Research from Start to Finish by Robert K. Yin. *International Statistical Review*, 79(3), 499–500. https://doi.org/10.1111/j.1751-5823.2011.00159_20.x
- Palley, T., Pérez Caldentey, E., & Vernengo, M. (Eds.). (2023). *Varieties of Capitalism*. <https://doi.org/10.4337/9781035312757>
- Rapacki, R., & Czerniak, A. (2018). Emerging models of patchwork capitalism in Central and Eastern Europe: empirical results of subspace clustering. *International Journal of Management and Economics*, 54(4), 251–268. <https://doi.org/10.2478/ijme-2018-0025>
- Roberts, H., Cows, J., Hine, E., Morley, J., Wang, V., Taddeo, M., & Floridi, L. (2022). Governing artificial intelligence in China and the European Union: Comparing aims and promoting ethical outcomes. *The Information Society*, 39(2), 79–97. <https://doi.org/10.1080/01972243.2022.2124565>
- Rolf, S. (2020, October 16). From Varieties of Capitalism to Uneven and Combined Development: A New Perspective. *China's Uneven and Combined Development*, 59–86. https://doi.org/10.1007/978-3-030-55559-7_3
- Rutten, Koen (2013) *Authoritarianism, capitalism and institutional interdependencies in the Chinese economy: implications for governance and innovation*. PhD thesis, London School of Economics and Political Science.
- Schmidt, Vivien A. (2007) *Bringing the State Back Into the Varieties of Capitalism And Discourse Back Into the Explanation of Change*. CES Germany & Europe Working Papers, no. 07.3., [Working Paper]
- Smuha, N. A. (2021). From a ‘race to AI’ to a ‘race to AI regulation’: regulatory competition for artificial intelligence. *Law, Innovation and Technology*, 13(1), 57–84. <https://doi.org/10.1080/17579961.2021.1898300>
- Turing, A. (2004). *Intelligent Machinery (1948)*. *The Essential Turing*. <https://doi.org/10.1093/oso/9780198250791.003.0016>
- The White House (2023). Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/>

Witt, M. A. (2010). China: What Variety of Capitalism? *SSRN Electronic Journal*.
<https://doi.org/10.2139/ssrn.1695940>

Yin, R. K. (2011). Qualitative Research from Start to Finish by Robert K. Yin. *International Statistical Review*, 79(3), 499–500. https://doi.org/10.1111/j.1751-5823.2011.00159_20.x

Zhang, C., & Lu, Y. (2021, September). Study on artificial intelligence: The state of the art and future prospects. *Journal of Industrial Information Integration*, 23, 100224.
<https://doi.org/10.1016/j.jii.2021.100224>

Zhang, J., & Peck, J. (2014, January 9). Variegated Capitalism, Chinese Style: Regional Models, Multi-scalar Constructions. *Regional Studies*, 50(1), 52–78.
<https://doi.org/10.1080/00343404.2013.856514>