



BUILDING A EUROPEAN AI ECOSYSTEM OF EXCELLENCE

Foresight Report

Susana Aires Gomes, Harry Crichton-Miller, Nicoleta
Kyosovska, Berta Mizsei, Davide Monaco, Laura Nurski,
Federico Plantera and Robert Praas

SUMMARY

We conducted a foresight exercise to explore how AI could reshape Europe's innovation systems and labour markets by 2045, for the CEPS project on a European Ecosystem of Excellence in AI. Through participatory scenario building and backcasting with 21 experts, the exercise linked global governance and R&D dynamics (macro level) with technological change in the workplace (micro level). Two integrated futures emerged. One depicted globally coordinated, market-driven, exponential growth and automation. The other envisaged globally fragmented, government-led, slower innovation, which is human-centred.

The process revealed that Europe's future competitiveness will depend less on technological speed than on institutional capacity – the ability to govern AI legitimately, distribute its gains fairly, and evaluate its social outcomes. Along both trajectories, effective procurement, worker participation, and evidence-based adaptation will be the foundations of an AI ecosystem that is both innovative and inclusive.



Susana Aires Gomes is a Researcher in the Global Governance, Regulation, Innovation and Digital Economy (GRID) unit at CEPS. Harry Crichton-Miller is a Research Assistant in the Economic Policy unit at CEPS. Nicoleta Kyosovska is a Research Assistant in the GRID unit at CEPS. Berta Mizsei is a Researcher in the GRID unit at CEPS. Davide Monaco is a Research Fellow in the Jobs and Skills unit at CEPS. Laura Nurski is a Research Fellow in the Jobs and Skills unit and Head of Programme on Future of Work at CEPS. Federico Plantera is a Researcher in the GRID unit at CEPS. Robert Praas is a Data Scientist at CEPS. The authors would like to thank Katja Spanz and Jaime Martinez for organisational support before and during the foresight workshop, Andrea Renda for his expert advice and guidance, and Daniel Cassidy-Deketelaere and Julien Libert for editorial support.

CEPS In-depth Analysis papers offer a deeper and more comprehensive overview of a wide range of key policy questions facing Europe. Unless otherwise indicated, the views expressed are attributable only to the authors in a personal capacity and not to any institution with which they are associated.

CONTENTS

1. INTRODUCTION	2
2. FORESIGHT METHODOLOGY	4
3. MACRO DRIVERS AND SCENARIOS.....	9
4. MICRO DRIVERS AND SCENARIOS	15
5. LINKING THE MICRO AND MACRO PLANES	20
6. BUNDLED FUTURE 1: THE HIGH-TECH, HIGH-DISPLACEMENT SCENARIO	23
7. BUNDLED FUTURE 2: THE SLOWER, MISSION-ORIENTED SCENARIO	30
8. CONCLUSION.....	37

FIGURES

FIGURE 1. VISUAL OF THE METHODOLOGY USED IN THE FORESIGHT EXERCISE	4
FIGURE 2. MACRO SCENARIOS	5
FIGURE 3. MICRO SCENARIOS	6
FIGURE 4. MACRO SCENARIO 1 BEFORE GROUP DISCUSSION	11
FIGURE 5. MACRO SCENARIO 2 BEFORE GROUP DISCUSSION	12
FIGURE 6. MACRO SCENARIO 3 BEFORE GROUP DISCUSSION	13
FIGURE 7. MACRO SCENARIO 4 BEFORE GROUP DISCUSSION	13
FIGURE 8. MICRO SCENARIO 1 BEFORE GROUP DISCUSSION	17
FIGURE 9. MICRO SCENARIO 2 BEFORE GROUP DISCUSSION	18
FIGURE 10. MICRO SCENARIO 3 BEFORE GROUP DISCUSSION	19
FIGURE 11. MICRO SCENARIO 4 BEFORE GROUP DISCUSSION	19

1. INTRODUCTION

This foresight exercise for the project on a European Ecosystem of Excellence in AI explores how AI could transform Europe's innovation landscape and labour markets by 2045. The objective, as always with foresight, is not to forecast a single future, but to map multiple plausible trajectories of technological, political, and social change. This enables us to discuss Europe's capacity to remain competitive, cohesive, and human-centred in the age of AI under alternative scenarios.

Foresight in this context serves a strategic function: it enables policymakers and stakeholders to expand their horizon and account for possible disruptive interactions between technology and society in the not-too-distant future. Rather than predicting outcomes, foresight identifies the boundaries of what might yet happen. Practices such as strategic foresight and backcasting then highlight and explore the choices that would make one trajectory (the 'preferred future') more likely than another. The resulting scenarios help reveal tensions, trade-offs, and opportunities that a policy-driven ecosystem for AI must address.

Bringing together 21 experts, the exercise builds on earlier CEPS research on AI governance, industrial policy, and the future of work. It extends this by examining how global dynamics and workplace realities interact. This integration of macro and micro perspectives bridges two debates that are often treated separately: one examining the geopolitical and economic architecture of AI, and the other delving into its effects on jobs and skills.

Europe's position in the global AI race makes this dual perspective essential. The region faces simultaneous pressures to catch up technologically and preserve social cohesion. On one hand, market forces led by private technology firms, many headquartered outside Europe, are pushing for rapid innovation and scale. On the other, Europe's regulatory and social models are prioritising safety, fairness, and human oversight. These opposing forces raise a core policy question: *can Europe design an AI ecosystem that is both globally competitive and socially sustainable?*

To test the limits of that question, the foresight process contrasted futures along two interlinked analytical levels. At the macro level, participants explored how AI might evolve under globally coordinated or fragmented governance, and whether research and innovation would primarily be market-driven or government-driven. At the micro level, they considered how quickly AI capabilities might progress on economically relevant tasks – following an exponential or logarithmic path – and whether these technologies would automate or augment human work. Together, these four axes formed a 'morphological

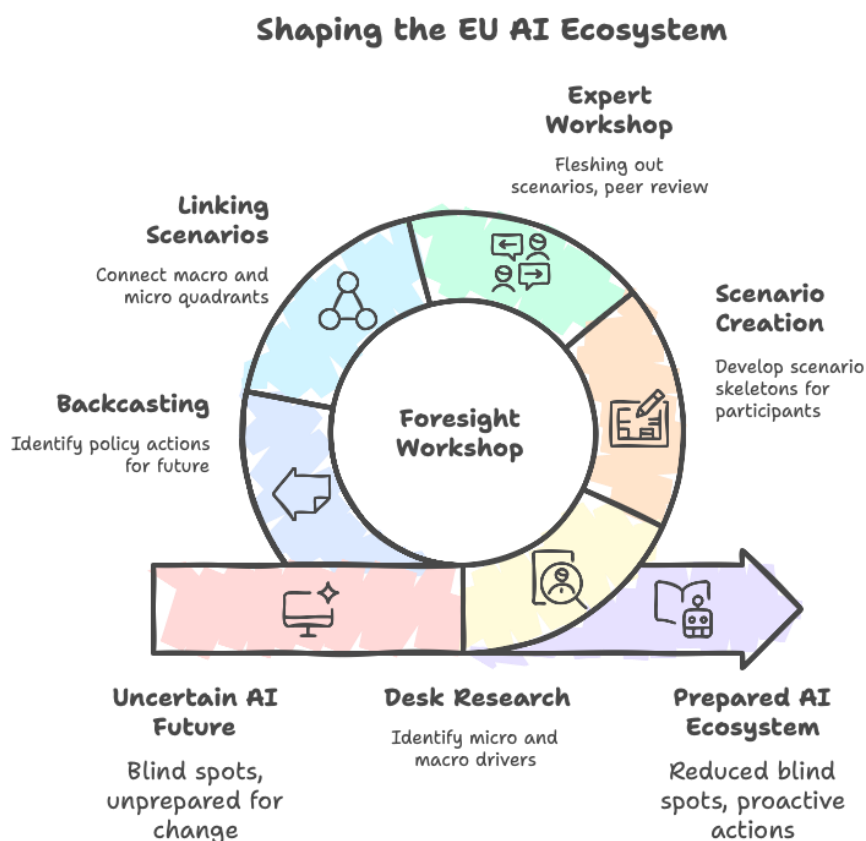
grid’ that structured discussions and enabled participants to investigate how combinations of political economy and technological change could shape Europe’s future AI landscape.

The outcome was not a set of predictions, but two internally coherent, bundled futures developed through participatory scenario building and backcasting. These scenarios envisage alternative pathways towards an AI ecosystem of excellence. They offer policymakers structured foresight on the conditions under which the EU might either strengthen or undermine its capacity to balance innovation with democratic and social resilience.

2. FORESIGHT METHODOLOGY

Foresight and futures thinking is not so much like peering into a crystal ball as it is looking into a kaleidoscope – instead of seeing one future, foresight encourages participants to observe multiple ones. Considering different futures, including fringe cases, in a structured way is meant to reduce blind spots, anticipate change and enable preparatory actions. In this case, CEPS worked with experts on AI and labour markets to conceive various future forms of an EU AI ecosystem. Figure 1 outlines the different steps taken throughout the preparation and a two-day workshop.

Figure 1. Visual of the methodology used in the foresight exercise



Made with Napkin

2.1. PREPARATION

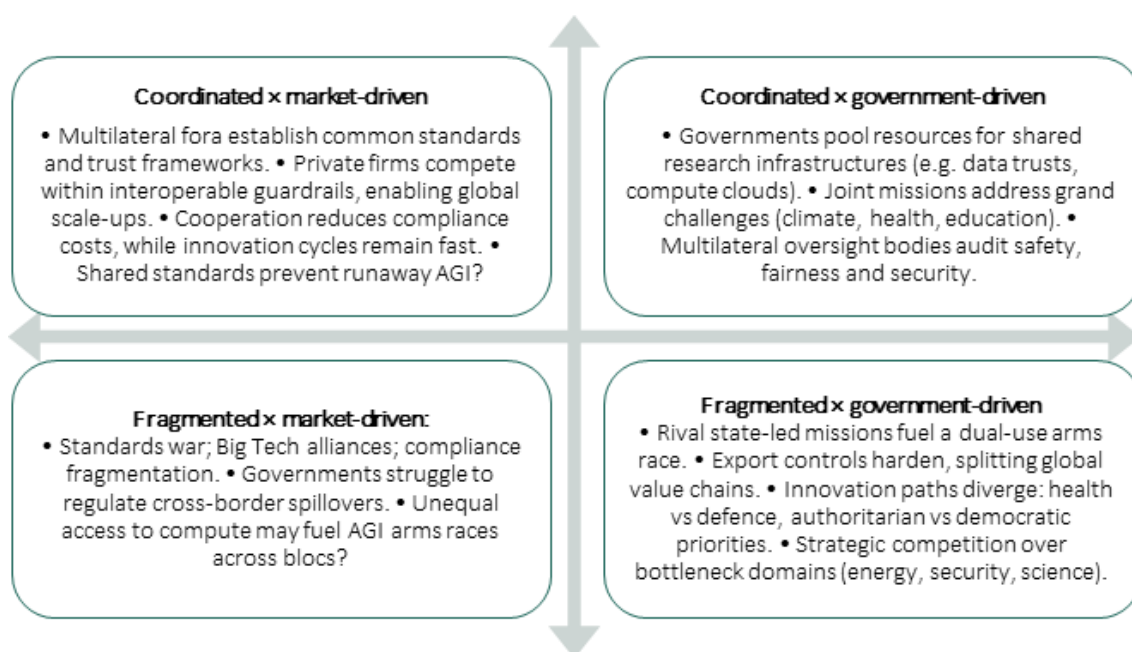
The foresight process started with desk research: CEPS researchers identified components that affect the EU's AI ecosystem and clustered these into micro and macro drivers. In foresight, drivers are issues, topics, trends, and other developments that affect change. Our micro drivers related to labour and technological development, while the macro drivers concerned social, economic, and political issues at the national and international

levels. Four key drivers were identified for the workshop: the two macro drivers were the structure of global AI governance and the nature of AI R&D, and the two micro drivers were the speed of development of AI capabilities and how AI changes work at the task level. These were then presented as four axes that make up two coordinate planes (one micro plane and one macro plane).

In the quadrants of the axes, CEPS researchers supplied scenario skeletons for the participants. The scenarios were images of possible futures to investigate in order to assess implications, monitor and prepare for threats, and develop ‘future-proof’ policies. They were not intended to be predictions or normative visions (i.e. the experts were not asked to comment on which future was most likely or most preferred). Both the drivers and the skeleton scenarios are presented in Sections 3 and 4 of this report.

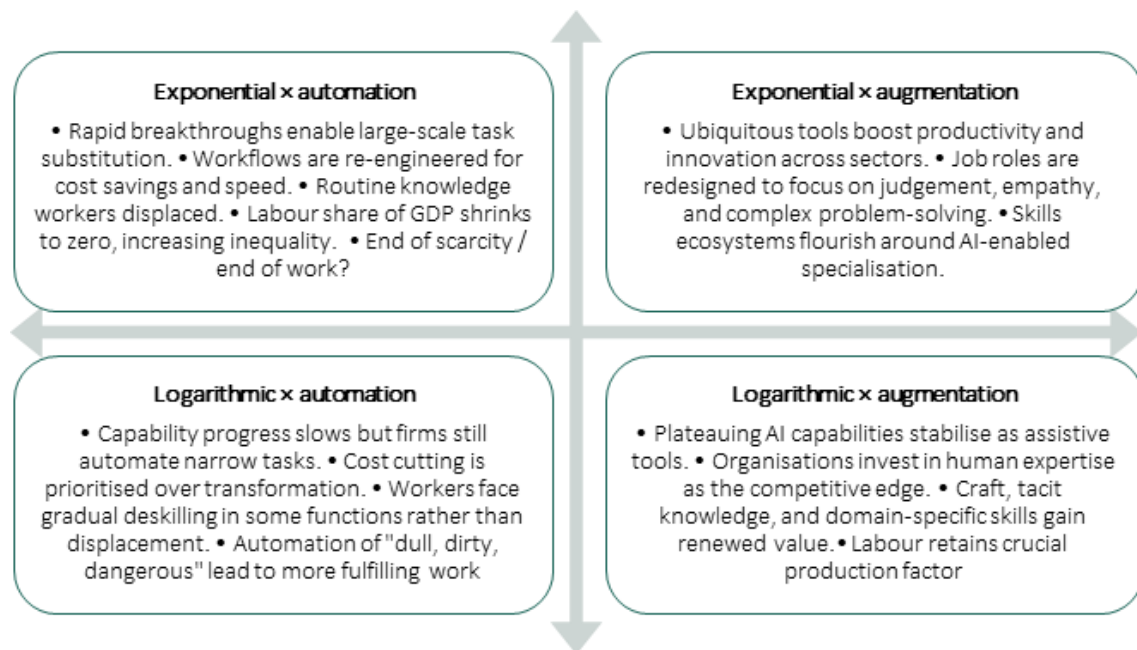
Expert participants with backgrounds relevant to the four key drivers were then identified and invited. The experts were sorted into micro-focused and macro-focused groups for the first day based on their personal preference, which were balanced for size and gender. The foresight workshop brought together 15 external experts and 6 CEPS experts, with a further 2 CEPS researchers moderating and 2 taking notes. More details on the participants can be found in Box 1.

Figure 2. Macro scenarios



Source: authors

Figure 3. Micro scenarios



Source: authors.

2.2. WORKSHOP FLOW ON DAY 1: SCENARIOS, PEER REVIEW AND LINKING

The first day of the workshop was dedicated to fleshing out the scenarios. One group worked on the macro plane and the other on the micro plane. Participants were asked to check the four scenario sketches provided to them and enrich them by thinking through what the drivers would look like in each quadrant. To aid this process and better situate themselves in the future, participants were asked to create ‘headlines from 2045’, prompted by moderators (e.g. *what do decision-making processes look like? What are the data, norms, and standards around AI?*).

Once the scenarios had been enriched, the two groups had the chance to peer review each other’s work for logical reasoning (i.e. *do the aspects of the scenario follow from the pre-assigned drivers?*). This was done through dot voting (blue dots for logically consistent aspects, red dots for perceived fallacies). Moderated discussions then considered any perceived contradictions. The original groups could subsequently decide on how to handle or incorporate the feedback. In multiple cases this led to a scenario being refined or substantially changed.

The first day ended with the expert participants (who by then had explored both the macro and micro planes) linking macro and micro quadrants by asking *which macro-micro quadrants were most likely to coincide*. Their votes and reasoning were noted and used

by CEPS researchers to pick macro-micro bundles for the second day's backcasting exercise. The team picked two bundles that maintained coherence in the micro-macro pairing according to the linkages exercise, but which were also sufficiently different from each other to result in a variety of policy approaches.

2.3. WORKSHOP FLOW ON DAY 2: BACKCASTING SELECTED BUNDLED FUTURES

For the second day, the groups were reshuffled to have a mix of micro- and macro-focused participants in both. Each group was then assigned one of the micro-macro bundles and asked what the EU's AI ecosystem might look like in that future. Finally, participants were asked to backcast policy actions that could make the most of the future bundle they had explored. In backcasting, experts move their way back to the present through a structured process. In this case, they identified the steps and conditions necessary to create an EU ecosystem of excellence in AI amid the constraints presented by the 2045 bundles.

Note that they were not working with a desired future, and it was conceptually difficult to think through policy actions that could strengthen the EU's AI ecosystem while remaining true to the constraints posed by the bundled futures (e.g. not having abundant resources due to protectionist trade blocs). To aid the process, participants were prompted to think through relevant questions for innovation and labour policy, such as:

- How to create a fertile environment for investment and capital?
- How to catalyse R&D and innovation around trustworthy, human-centric AI?
- How to ensure the availability of and access to compute infrastructure and data?
- How to retain, attract and train AI talent?
- How to ensure quality jobs both 'above' and 'below' the algorithm?
- How to avoid worsening inequality (between capital and labour, but also within the labour force and among capital owners)?

Sections 5, 6 and 7 of this report present the outcomes of the linking exercise, chosen bundles and backcasting.

Box 1. Participants in the foresight workshop

The foresight workshop brought together experts from research, policy, industry, and civil society to explore future pathways for a European ecosystem of excellence in AI. The list below includes participants who were present during the two-day workshop and consented to be named, as well as CEPS staff involved in facilitation, research, and documentation. Note: not all participants agree with all things mentioned in the report.

External participants

- Nathan Brandsen, IPSOS
- Andrea Glorioso
- Anna Milanez, OECD
- Pelin Özgül, ROA Maastricht University
- Alexander Petropoulos, Centre for Future Generations
- Sabine Richly, copyright & media lawyer
- Maud Sacquet and Catalina Gemanari, LinkedIn
- Siddhi Pal, interface
- David Timis, Generation
- Risto Uuk, Future of Life Institute
- Lieven Van Nieuwenhuyze, House of HR / World Employment Confederation
- Bart Van Rompaye, KPMG Belgium
- Three anonymous participants from the education, tech and public sector

CEPS participants

- Susana Aires Gomes
- Davide Monaco
- Laura Nurski
- Federico Plantera
- Andrea Renda
- Tim Schröder

CEPS facilitators and notetakers

- Facilitators: Berta Mizsei and Robert Praas
- Notetakers: Harry Crichton-Miller and Nicoleta Kyosovska

3. MACRO DRIVERS AND SCENARIOS

The macro plane explores how global dynamics shape the development and governance of AI through two key drivers: global governance of AI and AI research and development. These dimensions capture the geopolitical and institutional context that determines how innovation unfolds and who sets its priorities. The insights presented below draw from the scenario enrichment by the experts in the macro group, the peer review by the micro group, and the refinements made during the feedback round, as outlined in Section 2.

3.1. MACRO DRIVER 1: FRAGMENTED VS COORDINATED GLOBAL GOVERNANCE OF AI

Global governance of AI refers to the dynamics in international cooperation in setting priorities, policies, standards and safeguards for AI design, as well as its development, deployment and use in a variety of domains. If global governance is fragmented, this would entail lack of alignment on priorities and guardrails, and unilateral implementation of mechanisms for risk reduction. Excessive technological rivalry and divergence in policy approaches could become barriers to international cooperation. AI use is also increasingly intertwined with ethical and societal concerns which, if diverging, would present another obstacle to more cohesive global governance of AI.

Coordinated global governance of AI might imply multilateral fora that set global priorities for AI, interoperable standards for steering global AI development, and joint oversight mechanisms for monitoring AI use. Coordination might stem from alignment around AI risks or the need for coherent AI regulation for trade; but it could alternatively arise from the prevalence and dominance of certain non-democratic countries. Governance dynamics deeply affect key issues such as the agency and sovereignty of individual states, opportunities for innovation and trade in AI, and even progress on global challenges like climate change and public health.

3.2. MACRO DRIVER 2: MARKET-DRIVEN VS GOVERNMENT-DRIVEN AI R&D

The second macro driver considers who sets the direction of AI research and development. A market-driven path would be led by private capital and large technology companies, which could prioritise closed AI development to safeguard intellectual property and short-term goals such as profit. At the same time, given healthy competitive pressure, this path could enable quick AI development and high user satisfaction.

A government-driven trajectory could prioritise open source AI development, the availability of public infrastructure and the pursuit of long-term goals in the shape of mission-oriented research, focusing on sustainability and human flourishing. Nonetheless, the geopolitical context could push for more proprietary approaches, e.g.

for national security reasons, prioritising development for defence rather than global challenges. Similar to global governance, whoever leads AI R&D will influence the setting of AI priorities, and have a pronounced effect on the openness and pace of AI development and adoption.

3.3. FINDINGS FROM THE MACRO PLANE

Combining the two macro-level axes results in four possible configurations, yielding four distinct forecasting scenarios for the macro plane in 2045: (i) coordinated governance with market-driven R&D, (ii) fragmented governance with market-driven R&D, (iii) coordinated governance with government-driven R&D, and (iv) fragmented governance with government-driven R&D.

In the sections below, we provide an overall synthesis of the key patterns and insights emerging from the macro plane, followed by a detailed discussion of each scenario.

Participants engaged in lively and constructive discussions that reflected a high level of expertise and creativity. The most notable patterns and insights are summarised below.

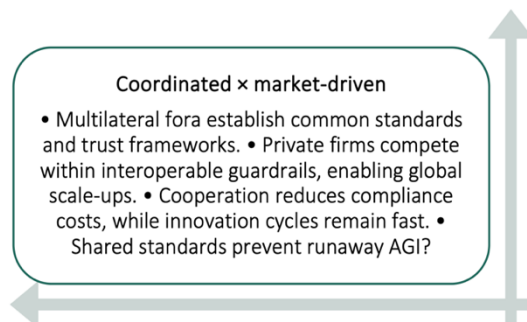
- The coordinated governance scenarios were found to have more potential for being positive, conditional on setting the right goals, with some participants suggesting the market-driven alternative as the preferred future.
- Many futures entailed a concentration of power into oligopolies or technocracies, the former belonging to the market-driven R&D quadrants and the latter to government-driven R&D.
- Discussions on the effects of fragmentation in governance revealed a duality: uncertain enforcement of cross-border standards vs autonomy and within-border innovation.
- Global conflicts dominated the government-driven R&D scenarios, with a slightly stronger presence in the fragmented governance version.
- The public had a weak role in most scenarios, with a slightly more important role in market-driven R&D but eroded by the nudging of user preferences.
- Experts had an overall negative outlook on the prospects for progress on global challenges and societal values driving AI development.
- The future role of the EU was a contentious topic: would it lag further behind and become irrelevant or would its supranational coordination become a model for other states?

The findings for each scenario are presented in detail in the next sections.

3.4. SCENARIO 1: COORDINATED GLOBAL GOVERNANCE X MARKET-DRIVEN R&D

The first scenario the group considered was the top-left quadrant – coordinated global governance of AI and market-driven AI R&D. Based on the directions provided, participants expected fast innovation but with the main drivers being profit and user feedback, rather than challenges such as climate change. They initially considered weak state power, with standards being tech-led and implementation remaining uncertain. One participant characterised multilateral fora as ‘expos’ where states compete for corporate attention.

Figure 4. Macro scenario 1 before group discussion



Source: authors

Challenged by peer review, an alternative, more positive future emerged in which coordinated governance empowers governments and leads to regulated AI development and proper enforcement of trust frameworks. In this respect, weak state power would most likely to lead to self-regulatory solutions, which could produce good results if civil society or governments are empowered to peer

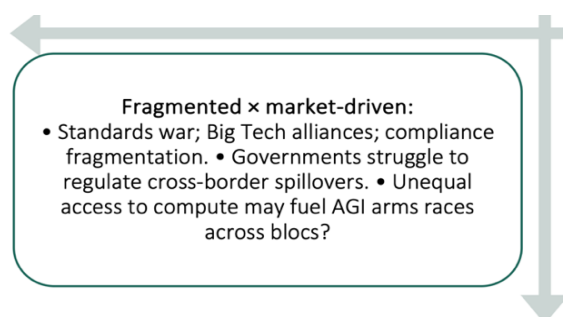
review and audit the AI systems that private players deploy on the market.

Another fork in the quadrant's trajectory, conditional on government power, concerned the presence of big tech and oligopolies. Concentration of market power could either remain high, so far as to take the form of ‘united companies’, or disperse, giving way to fair and functional markets.

3.5. SCENARIO 2: FRAGMENTED GLOBAL GOVERNANCE X MARKET-DRIVEN R&D

The second scenario also concerned market-driven AI R&D, but it was paired with fragmented global governance – the bottom-left quadrant. Many participants mentioned large-scale conflicts: either in the form of trade (tariffs, export controls, or restrictions on FDI) and regulation wars, or physical war due to AI being used by the military without constraints. The group envisaged geopolitical tension over raw materials and energy, and the possibility of (neo)neocolonialism, although peer reviewers noted that this could be a more pronounced feature of government-driven AI R&D.

Figure 5. Macro scenario 2 before group discussion



Source: authors

Visions around the balance of power between industry and the state were dispersed. Some felt that companies would be fully-fledged geopolitical actors (what is sometimes called a ‘techno-polar world’). Others conceived of states having the important role of mediators between companies, or imagined that states with big tech would be the most powerful actors – although the market-driven nature of this scenario might weaken the

connection between states and AI.

On the role of the public, there was agreement that more trust would be placed in companies than in countries. One participant questioned future user relevance, given that big tech could advertently and successfully shape user opinion.

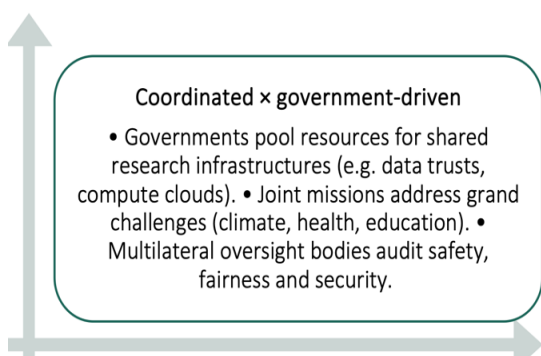
The idea of ‘CEO leaders’ (i.e. a situation in which AI commercial and technological power coincides with political leadership) was put forward. Yet some of the participants doubted the readiness and fitness of tech leaders for politics, while others challenged future, heavily media-influenced perceptions of leadership.

3.6. SCENARIO 3: COORDINATED GLOBAL GOVERNANCE X GOVERNMENT-DRIVEN R&D

Moving towards futures with government-driven AI R&D, the group explored the territories of the top-right quadrant, which adds coordinated global governance. Participants pictured a mostly negative future with state control giving rise to rather dystopian innovations (e.g. implanted chips embedding citizenship IDs) and public-private partnerships on efforts to link social behaviour with credit systems (social credit scoring).

Peer reviewers felt these were not a distinctive feature of the scenario, but the group deemed it a likely outcome of powerful states equipped with advanced technology.

Figure 6. Macro scenario 3 before group discussion



Source: authors

Research and development were portrayed as slow and bureaucratic, mission-driven or wasteful, or something in between with spill-over effects: ‘Crazy scientists waste taxes trying to teach AI how to juggle. Also, we cure cancer.’ The exclusivity of state-controlled research was illustrated by restricted youth programmes and crackdowns on rogue AI hacker fairs. One participant highlighted the paradox between the need for state control and the benefits of researcher freedom for

innovation.

The peer review process found gaps in the discussion, like the absence of non-state actors or considerations of data ownership and global voting power. Participants took note and re-considered industry as a participant in the state-led R&D.

3.7. SCENARIO 4: FRAGMENTED GLOBAL GOVERNANCE X GOVERNMENT-DRIVEN R&D

Finally, the workshop turned to the bottom-right quadrant of government-driven R&D and fragmented governance. This scenario took the shape of a fragmented world of sovereign

Figure 7. Macro scenario 4 before group discussion



Source: authors

AI, where states pursue self-reliant AI ecosystems reminiscent of a space race, going so far as to destroy physical cross-border infrastructure like transatlantic data cables. These global blocs were further imagined as utilising different forms of AI – for example, social credit scoring would be a feature of an authoritarian bloc. One participant suggested an AI winter, in which public investment bails out on European AI SMEs. Defence expenditure

would see an all-time high with some envisaging global war and declining social well-being.

Another participant's ingenuity turned the EU into a theme park as the only way for it to stay relevant alongside the US and China. Some welcomed this vision, arguing that the EU's weakness was a simple extrapolation of the present while others noted it was not a distinctive feature of this scenario.

AI stagnation was another contentious aspect: some felt fragmentation could instead foster within-border innovation enabled by the ability to resist dominant powers. The concluding remarks opened the possibility for a more positive scenario with states enjoying strategic autonomy and safeguarding democratic rights.

4. MICRO DRIVERS AND SCENARIOS

The micro plane examines how technological progress and workplace transformation interact through two main drivers: the pace of AI capability development and the nature of AI's impact on the workplace. These dimensions reveal whether AI evolves rapidly or plateaus, and whether it substitutes or complements human labour. The findings summarised below are based on the scenario building work of the experts in the micro group, the peer review by the macro group and subsequent adjustments made after group reflection, as described in Section 2.

4.1. MICRO DRIVER 1: EXPONENTIAL VS LOGARITHMIC DEVELOPMENT OF AI

On the first axis, AI capabilities might continue to accelerate towards artificial general intelligence (AGI) because of vast new resources – data, capital and energy – fuelling improvements, supportive policies fostering innovation, and widespread integration of AI into economies. If so, breakthroughs would continue to scale rapidly, and, by 2045, we could see systems surpassing human-level performance across an expanding range of cognitive tasks, fundamentally reshaping industries, scientific discovery, and even social structures.

Alternatively, progress could plateau if we face diminishing returns from larger models and low integration of AI, if public trust falls and regulation tightens around safety, or if we encounter data, physical hardware, capital or energy constraints that slow further scaling. If so, capabilities would stabilise around where they are today (or slightly above), with systems functioning as effective assistive tools capable of performing narrow tasks but not without supervision of a human, and human labour retaining a competitive edge.

4.2. MICRO DRIVER 2: AUTOMATION VS AUGMENTATION OF HUMAN CAPABILITIES

On the second axis, AI might change the nature of work in two directions. In one direction, via automation, machines would substitute large parts of human labour, especially in routinised or measurable tasks. While this could increase efficiency and productivity for some, by 2045, it might have displaced labour and caused aggregate unemployment, driving tensions between the 'haves' and 'have nots'. This might be because AI continues to dramatically improve, because the technology diffuses into businesses and the public sector and more effective use-cases are found, or because economic incentives push cost-cutting over workforce development.

In another direction, AI might change the nature of work via augmentation, where human-AI collaboration boosts the capacity of workers, enhancing judgement, creativity or 'soft' skills, to which technology serves as a complement. By 2045, new types or styles of jobs might have evolved that embed and effectively utilise technology, but this

transition could also exacerbate inequalities and require reskilling across the workforce. This might be driven either by technological or organisational design, or by regulation that fails to protect workers from, or adapt them to, the changing nature of work.

4.3. FINDINGS FROM THE MICRO PLANE SCENARIOS

Combining the above two axes gives us four possible combinations, in turn creating four possible forecasting scenarios for the micro plane in 2045: (i) logarithmic-augmentation, (ii) exponential-automation, (iii) logarithmic-automation, and (iv) exponential-automation. The following sections give an overall summary, drawing out themes and commonalities on the micro plane, before presenting each scenario in detail.

Extensive, productive debate among participants and via peer review led to four vivid, divergent scenarios. Nonetheless, an important consensus emerged on overarching themes, with the following key points.

- The ‘limitations’ in AI capability development would not just be regulatory, but also related to the availability of technology, as well as financial and natural resources. Trust, power, and ownership would also constrain AI adoption, meaning that analysing AI’s impact through a task-based or technical feasibility lens is too limited.
- All scenarios projected an overall drop in labour demand, even scenarios with a slower, logarithmic rate of AI development. The distinction between augmenting and automating AI was not always clear-cut, as in practice even ‘augmenting’ systems involve the ‘automation’ of some tasks.
- Job content would change dramatically across all scenarios – particularly the exponential ones – with some participants doubting whether labour markets would remain similar to what they are today. No scenario predicted automatic job-quality gains; improvements were seen as dependent on government or organisational action.
- Demographics were seen as a pivotal factor in all scenarios. If demographic collapse reduced the workforce, labour would gain relative value, job quality could rise, and social protection systems would face less strain.
- Numerous questions arose about the governance of AI and society, especially in the scenarios with exponential technological development. Key issues would include who owns productive technologies, who designs protection systems or regulation, and how much influence public opinion holds. Similarly, all scenarios

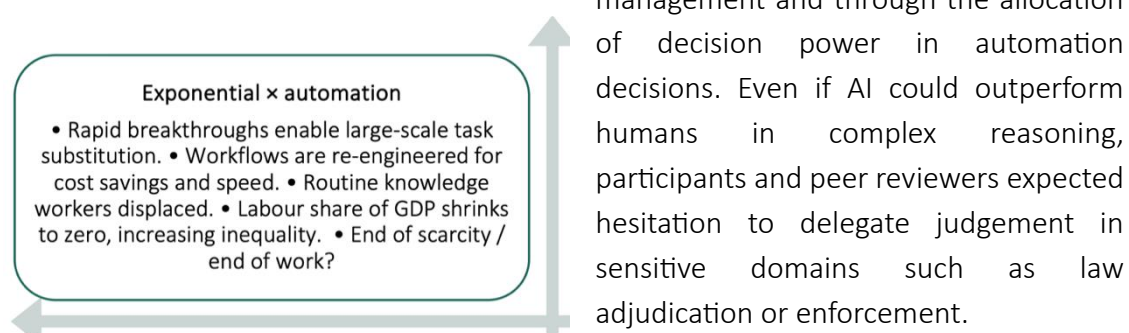
envisaged at least some rise in inequality. The main actors to stem this tide included unions, governments and businesses themselves. Participants reaffirmed the centrality of human agency in designing all four scenarios, summarised by one as ‘the future is not a prediction problem but a design problem’.

We present the findings for each scenario in detail below.

4.4. SCENARIO 1: EXPONENTIAL AI DEVELOPMENT X AUTOMATION OF WORK

Participants generally found this to be the most extreme scenario, though they stopped short of predicting ‘the end of work’. They expected humans to continue preferring services from other humans, and some workers, especially with relative power (e.g. doctors in hospitals) to block automation. Participants reasoned that the human desire for meaning would lead to prioritising human activity in areas people naturally find meaningful, such as caretaking and education. This insight suggests that the impact on job quality might be underappreciated compared with job displacement.

Figure 8. Micro scenario 1 before group discussion



Source: authors

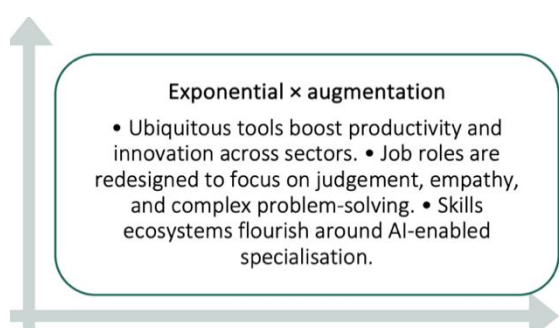
At the same time, economic concentration and declining labour share were seen as likely outcomes. Participants debated whether a few large AI firms would dominate or whether specialised companies might thrive through domain expertise. Either way, a redistribution of income and resources would become central: if earnings for workers shrink and profits concentrate, would society act quickly enough to establish a universal basic income (UBI) or windfall taxation? The group anticipated political tension, reform, and regulation in a post-labour economy.

4.5. SCENARIO 2: EXPONENTIAL AI DEVELOPMENT X AUGMENTATION OF WORK

This was generally viewed as the most optimistic scenario, with exponential progress driving a technological revolution comparable to that of the steam engine, complete with upheaval, job destruction, and net job creation and transformation.

Some envisaged a world where individuals become CEO of their own ‘one-person complex company’, with boosts to job quality, autonomy and productivity as AI handles complex marketing, administration and production functions. Peer review questioned the scalability of this model, noting that vulnerable groups might struggle to compete and that many people might not desire constant complexity in their work. A central tension emerged between productivity gains and job intensity: would AI reduce workload or fuel algorithmic optimisation and stress, reviving the ‘Luddite scenario’ of life under the machine’s rhythm?

Figure 9. Micro scenario 2 before group discussion



Source: authors

A key issue in this scenario would be ensuring that these transformations do not deepen inequality. Not everyone – for instance, manual workers – has skills that AI can effectively augment, and most benefits would accrue to (tech-owning) capital rather than labour. Further, participants worried about market concentration, as an exponential AI market – a general-purpose technology *prima facie* valuable in most sectors – would presumably be a ‘winner takes all’

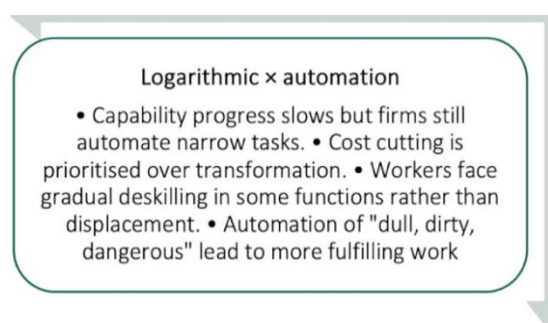
one.

This prospect raises critical governance and redistribution questions about ownership and control of technology. Stronger unions and collective institutions were seen as necessary to secure quality work, counterbalance capital, and ensure that AI truly augments rather than displaces human capabilities.

4.6. SCENARIO 3: LOGARITHMIC AI DEVELOPMENT X AUTOMATION OF WORK

There was a strong consensus among participants that the assumption that AI will primarily replace ‘dull, dangerous and dirty’ work was inaccurate. Much of this work remains resistant to automation, while many of the tasks most at risk – such as translation or communication – are safe, creative, and even pleasurable. On a logarithmic trajectory, this existing pattern of automation was expected to persist, but flatten. Peer reviewers cautioned that trends draw on short-term data and that labour market effects unfold slowly. One example was that of the slowed hiring of junior roles, which is part of a broader economic and organisational trend that, nonetheless, AI may have helped accelerate.

Figure 10. Micro scenario 3 before group discussion



Source: authors

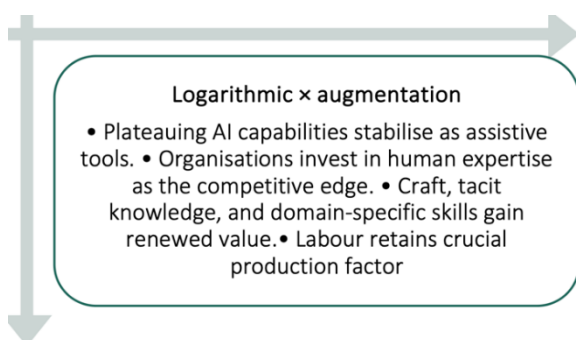
This scenario puts social protection systems under pressure, exposing the long-standing tension between Europe's welfare commitments and the need for competitiveness and growth. Participants noted that – in this scenario more than others – the depth of deskilling and displacement would largely depend on policy choices rather than on technological inevitability. The group warned of a 'sleepwalking scenario', in which gradual change dulls political attention until the

consequences become entrenched. Without deliberate policy intervention to steer AI use towards inclusion and resilience, policymakers (including EU ones) might find themselves balancing precariously between stagnation and decline.

4.7. SCENARIO 4: LOGARITHMIC AI DEVELOPMENT X AUGMENTATION OF WORK

Even modest, logarithmic and augmentative progress in AI was seen as sufficiently disruptive to reshape work and production in Europe by 2045. Participants compared this with the development of cars – steady improvements without transformation – but agreed that even such incremental change could have significant cumulative effects

Figure 11. Micro scenario 4 before group discussion



Source: authors

across sectors. The slower pace was perceived positively: it would give workers and firms more time to adapt, retrain, and redesign work. Yet this also posed policy dilemmas: should support focus on the already displaced, or on those still at risk of future disruption?

Participants anticipated that workers 'above the algorithm' – those in high-skilled creative or analytical roles – would benefit most, while mid-skilled professionals could lose ground as their

work becomes more easily replicated. By contrast, manual and low-skilled labour might gain relative value if non-automatable physical work becomes scarcer. The result was a gradual, rebalanced labour market, not free of inequality but potentially more adaptable. This scenario captured a world of managed transition, where social stability depends on steady investment in learning systems and inclusive technology governance.

5. LINKING THE MICRO AND MACRO PLANES

The foresight process treated the macro and micro levels as distinct analytical levels of a larger interdependent system shaping the European AI landscape. The end of day 1 of the workshop was structured to make this connection explicit. After developing four scenarios at each level – two macro (global governance and R&D orientation) and two micro (technological trajectory and workplace effects) – participants reconvened to identify which combinations were *most likely to co-evolve*. The aim was not to produce all sixteen theoretical combinations, but to reveal the underlying causal logic connecting international governance models, innovation dynamics, and transformations in work.

5.1. LINKING METHOD

In the afternoon ‘linking’ exercise, each group reviewed the four scenarios from the other level and answered two symmetrical questions:

- (i) *What is the most likely micro scenario following from each macro one?*
- (ii) *What is the most likely macro scenario underlying or generating each micro one?*

Participants first reflected individually, then discussed in pairs and small groups before voting. The process was designed to elicit causal plausibility rather than preference – to trace which futures seemed internally coherent, not which were desirable. The moderators recorded the links by tallying votes on flipcharts without debate, producing a transparent picture of convergence.

5.2. OBSERVED CONVERGENCE

The results showed a strong and consistent pattern. Both macro groups associated *market-driven* R&D with *exponential* technological development and *automation* of tasks, while *government-driven* R&D was overwhelmingly linked to *logarithmic* development and *augmentation* of human capabilities. In other words, the two micro dimensions effectively *collapsed* onto one of the macro axes. The caveats to this are discussed below.

Participants reasoned that when innovation is dominated by private capital and competition, speed and scale become decisive. This market-led environment favours rapid, exponential progress and cost-driven substitution of labour – automation being the natural outcome of relentless efficiency pressures. Firms such as OpenAI explicitly define their mission as developing ‘highly autonomous systems that outperform humans at most economically valuable work’ in the pursuit of AGI¹. These ambitions reinforce the

¹ See the [OpenAI Charter](#).

perception that market-driven R&D naturally aligns with exponential progress and widespread automation. By contrast, when governments steer R&D through mission-oriented programmes, they aim to avoid social disruption and impose constraints that slow development, emphasise public value, and embed human oversight. The result is a more gradual, logarithmic trajectory focused on augmenting rather than replacing workers.

The consensus was thus grounded in causal structure, not probability or normative choice. Participants saw macro conditions as determining the parameters of micro behaviour. More specifically, who funds AI research and for what purpose shapes both the speed at which capabilities grow and how they are applied in workplaces. The exercise produced narratives of *co-evolution*, where the institutional drivers of innovation cascade down to affect organisational design and job content. The aggregate pattern therefore aligned with two dominant macro-micro causal chains:

- **Chain 1** – Market → Exponential & Automation
- **Chain 2** – Government → Logarithmic & Augmentation

5.3. DIVERGENCES AND ALTERNATIVE VIEWS

The observed convergence above seems to take for granted that market players are capable of achieving AGI. However, it is on the cards that AI development hits a winter period regardless of the freedom and funding granted to market players.

Furthermore, discussions among the participants revealed several important nuances and alternative logics. A minority argued that exponential progress could also occur under government-driven R&D if states pooled resources efficiently, avoided political inertia, and adopted mission-driven innovation models. Some suggested that coordinated global governance might accelerate markets rather than slow them, by lowering compliance costs and harmonising standards. Others contested the assumption that fragmentation necessarily leads to dysfunction, noting that regulatory diversity could also generate innovation niches and resilience.

These debates underscored that the relationships between governance, innovation, and work are multi-causal rather than linear. The overall direction of agreement, however, remained clear even as participants recognised legitimate alternative logics.

5.4. ALLOCATING GOVERNANCE AND CREATING BUNDLES

With the main causal chains established, the organising team met after day 1 to determine how to pair the final remaining driver, namely the governance dimension (fragmented vs coordinated), with each chain. The straightforward assignment –

fragmentation with the ‘doom’ trajectory and coordination with the ‘resilient’ one – was deliberately inverted. The team chose to pair coordinated governance with the market-exponential-automation chain and fragmented governance with the government-logarithmic-augmentation chain. This inversion produced two more nuanced and policy-relevant futures:

- **CMEA** – coordinated governance, market-driven R&D, exponential development, and automation, representing a world of global standards and cross-border markets generating technological abundance yet also large-scale labour displacement; and
- **FGLA** – fragmented governance, government-driven R&D, logarithmic development, and augmentation, which is a slower, inward-looking, human-centred future where national strategies diverge yet social cohesion remains strong.

The decision aimed to avoid caricatured outcomes – an unmitigated dystopia or utopia – and instead generate tension within each bundle. In CMEA, coordination prevents the ‘AI singularity’ and systemic collapse but cannot protect all workers from displacement. In FGLA, fragmentation constrains utopian, global cooperation but allows plural experiments in inclusive, mission-driven innovation.

By the close of day 1, the workshop had thus extracted two integrated storylines from the 4×4 morphological grid of all possible macro-micro combinations. These insights provided the basis for the backcasting sessions on day 2, where groups explored what a European ecosystem of excellence in AI would require under each bundled future.

6. BUNDLED FUTURE 1: THE HIGH-TECH, HIGH-DISPLACEMENT SCENARIO

Following the first day of foresight, a group of participants completed a backcasting exercise on CMEA, which combines the macro scenario of *coordinated global governance* and *market-driven R&D* with the micro scenario of *exponential development* and *automation*. The aim was to identify the concrete policy actions and priorities needed to achieve the outcomes envisaged during the scenario work. Before the discussion, the bundle was presented to participants as follows:

The CMEA bundle depicts a coordinated yet highly technological world, marked by exponential growth in AI capabilities and market-led innovation, alongside a turbulent labour market marked by rapid displacement. The innovation ecosystem is profit-driven, delivering scalable breakthroughs, while global coordination through multilateral fora establishes safety and interoperability standards as guardrails on AGI. Despite the global coordination to contain uncontrolled AGI, major societal transformations occur. Large-scale automation of routine and cognitive tasks restructures workflows for efficiency and cost reduction. This has a major impact on workers who face a highly competitive and volatile labour market. Policymakers attempt to protect people through urgent reforms but struggle to keep pace with the speed of change.

6.1. WORKSHOP JOURNEY

The backcasting exercise for bundle 1 began with world-building: participants had to stabilise the scenario before policy work could begin. The CMEA bundle presented internal tensions that required explicit negotiation. Was this a resource-rich future, perhaps even approaching post-scarcity, or would resource constraints persist? Would the EU still exist as a coherent political entity in 2045, or would coordination have shifted to other scales?

These questions revealed a deeper challenge – exponential technological development combined with coordinated governance produces futures that strain conventional economic logic. If AI capabilities grow exponentially while coordination prevents systemic collapse, what does the resulting political economy actually look like? Participants had to construct shared assumptions about resource availability, institutional continuity, and the basic mechanics of wealth generation before they could meaningfully discuss policy.

Once a baseline scenario was set, they explored policy priorities across six thematic categories: capital, R&D, data and compute, talent, job quality, and inequality. Moderators

used these as prompts to identify goals and policy tools for a European ecosystem of excellence under the CMEA constraints. Three moments of analytical friction recurred across the discussion.

First, participants encountered what might be called the economic incoherence problem. If AI makes production extremely efficient, driving costs towards zero, what happens to capital returns? Who invests, and why, when traditional profit mechanisms collapse? As one participant asked, only half-joking: ‘If AI makes everything cheap, does everyone become rich or bored?’ This pointed to a genuine puzzle: an economy with abundant productive capacity, but concentrated ownership of AI infrastructure, produces outcomes that resist easy categorisation. References to the film *Elysium* functioned as shorthand for a two-class system where technological abundance coexists with severe stratification, but participants could not fully resolve how such a system would remain economically or politically stable².

Second, the redistribution impasse surfaced across multiple topics. What would actually be redistributed – money, compute resources, personal AI agents, or robot ownership? Should redistribution occur on a national, European or global scale, and through what mechanisms? Participants discussed windfall taxes on AI profits, UBI funded by robot taxation, and public investment in AI companies through pension funds. Yet each proposal triggered immediate concerns about legitimacy, implementation, and unintended consequences. As one participant noted, ‘governments might be less trusted than companies’ (especially in an AGI age), raising the question of who could credibly administer large-scale redistribution when the traditional institutions of social democracy face eroded public confidence.

Third, participants struggled with the timeline problem. Would exponential AI development and the resulting workforce displacement occur faster than policy responses could form? Some argued that redistribution mechanisms could be implemented rapidly, ‘like the New Deal’, given the political will to do so. Others questioned whether democratic systems, operating on electoral cycles and constrained by legacy institutions, could match the pace of market-driven automation. The possibility of riots, mass unemployment, and social unrest arose in the discussion not as alarmist speculation but as realistic governance challenges if displacement significantly outpaced redistribution.

These frictions were not resolved but served as pressure points shaping where future fault lines in policy might emerge. As the exercise shifted to backcasting from 2045,

² Note that the logic of the bundle would also allow for costs to move upstream to AI factories.

participants moved from describing ideal policy landscapes to confronting the political difficulty of implementation. The Capital Markets Union, for instance, was acknowledged as essential for enabling start-ups to operate across Member States but it was also recognised as ‘long-term, politically difficult, and not easily made feasible’. The establishment of a proposed 28th *regime* for companies could be a starting point in that direction. But proposals for institutional reform – such as allowing civil servants to use more complex AI tools or creating EU-level governance structures – immediately triggered discussion of sovereignty concerns and bureaucratic resistance.

The backcasting process revealed a gap between normative ambition and political realism that participants could not fully bridge. Ideas sorted themselves into categories of feasibility:

- relatively easy and short-term measures (tax incentives for company investment in reskilling and sectoral data quality standards),
- harder short-term tasks (worker representation across professions and performance measures that include well-being alongside productivity), and
- hard long-term goals (redistribution of productivity gains, active citizenship models, and reduced working hours).

This *taxonomy of difficulty* brought to the surface the core dilemma of foresight work in the context of exponential change and coordinated governance: the policies that seem most necessary are often those least compatible with existing institutional capacity.

6.2. POLICY ACTION AND UNDERLYING DYNAMICS

Across the six thematic prompts, the conversation coalesced around five underlying dynamics that revealed the underlying architecture of how participants conceived the CMEA future.

Geography and scale came up throughout the discussion but never stabilised as a clear policy logic. Participants could not agree on whether an AI-driven economy would centralise or decentralise human activity. The question surfaced in debates about talent policy (should Europe focus on retaining its own researchers or attracting global talent?), compute infrastructure (concentrate resources in designated zones or distribute them more broadly?), and inequality (will disparities grow between urban knowledge centres and rural regions?). Some participants argued that AI reduces the need for physical clustering in cities, allowing people to live anywhere (depending on connectivity) while remaining economically productive. Others countered that cities would retain importance as sites of cultural experience and human encounter, even if traditional work no longer

requires co-location. The debate remained unresolved, pointing to deeper uncertainty about the spatial organisation of life and work in an AI-intensive economy.

Redistribution functioned as the structural question cutting across multiple policy domains rather than a discrete issue contained within the inequality category. It arose in discussions on capital (who owns or should own the means of production – chips, data centres, and foundational models), R&D (who sets priorities for AI innovation and reaps its benefits?), job quality (how are gains shared between capital and labour?), and dealing with inequality itself (through which mechanisms – UBI, taxation or public investment – and at what scale?). Participants ultimately understood redistribution not as a single policy tool but as the central governance problem of the CMEA scenario. If technological capability grows exponentially while ownership remains concentrated, redistribution will become the mechanism through which coordination either maintains social cohesion or fails to do so.

Trust and legitimacy emerged as a constraint on governance capacity that participants considered more binding than technical feasibility. The issue cropped up in discussions of who should administer redistribution (governments or companies), who should set R&D priorities (the scientific community, public panels, market forces, or some combination of these), who should govern data quality (independent stewards, sectoral bodies, or state regulators), and how to maintain public confidence in both political and technological elites. In a ‘machine-dependent society’, how does public trust in governance evolve when technological elites operate with greater perceived competence and responsiveness than political institutions? This reflects concern that democratic systems might lack the legitimacy needed to implement the scale of redistribution and institutional reform that the CMEA scenario would require. Furthermore, in a globally coordinated scenario with AGI, standards might be driven by dominant market players rather than governments.

Positioning ‘above’ or ‘below’ the algorithm thresholds structured thinking about inequality in ways that extended beyond employment status. The phrase initially referred to job quality, distinguishing those who build or deploy AI from those whose work is directed by AI systems. But the concept migrated across multiple domains. In talent policy, it shaped discussion of who would be ‘builders’ versus ‘users’ of AI capabilities. In inequality discussions, it surfaced as ‘tech poverty’ – the possibility that those without access to personal AI agents would experience a new form of deprivation. In civic participation debates, it appeared as the difference between those who set tasks and those who perform them. Participants ultimately understood the CMEA future as

producing not just wealth inequality but positional inequality: disparities in autonomy, agency, and the capacity to shape one's relationship with automated systems.

Feasibility and political time became explicit during the backcasting exercise, when participants had to assess which policies were achievable given the constraints. A taxonomy of difficulty revealed implicit assumptions about institutional capacity. Some measures were deemed relatively easy: tax incentives for firm-offered training, public-private partnerships for compute infrastructure, and industry-led, sectoral standards on data quality. Others were acknowledged as harder: ensuring worker representation across automatable professions, shifting performance measures from productivity to well-being, and enabling rapid AI deployment in the public sector. Long-term and politically difficult measures included fundamental reforms at the EU level, like the Capital Markets Union, large-scale redistribution mechanisms, and the redesign of work around active citizenship rather than employment. This sorting revealed a mismatch between the ambition required by the scenario and the institutional capacity of political systems to deliver transformation fast enough to match exponential change.

6.3. KEY TAKEAWAYS

The bundle 1 backcasting exercise brought up several insights about the conceptual and practical limits of policy foresight when confronting exponential technological change.

Automation remained conceptually unstable. Participants could not fully resolve the economic logic of a future in which AI handles 'most economically valuable work' while market-driven innovation continues. Traditional mechanisms of capital accumulation assume scarcity and the need for human labour to convert inputs into outputs. If AI dramatically reduces both, what sustains investment, profit, and growth? The references to post-scarcity, zero-cost production and concentrated ownership of AI infrastructure pointed towards a future that breaks conventional economic assumptions. The *Elysium* metaphor – a two-class system of abundance and deprivation – became shorthand for contradictions the group could not resolve. As one participant noted, such a system would 'not [be] pleasant'. It would also potentially be unstable, because 'if only 1% control *everything*, the system collapses'. What might prevent or follow that collapse remained open.

Universal basic income functioned as conceptual scaffolding rather than a coherent solution. UBI was continually mentioned throughout the discussion – funded by robot taxes, windfall clauses, energy taxation, or data transaction fees; administered on a national, EU, or global scale; tied to civic participation or provided unconditionally. Yet it never developed into a fully specified policy mechanism. Participants questioned whether UBI represented charity or justice, empowerment or dependency, who would administer

it and how much would be enough. These questions were not answered definitively. UBI served as a conceptual anchor for discussions of redistribution without solving the underlying problems of legitimacy, motivation, and political feasibility. Its recurrence suggests that it functions more as a symptom of the redistribution impasse than as a resolution and that it demands its own fully-fledged foresight workshop to make it future-ready.

The collapse of work raised questions of purpose that policy categories could not easily accommodate. If AI automates most economically valuable tasks, what organises human life? Money ceases to function as the primary incentive for education, research, or civic engagement. Participants noted that ‘humans naturally recreate organisational forms’ even in the absence of traditional work but struggled to articulate what forms these would take. Proposals for active citizenship models – where UBI is tied to participation in local councils or community projects – immediately triggered concerns about social credit systems and paternalism. The suggestion that future inequality might be determined by access to ‘monetised companionship’ or status in online social hierarchies pointed towards a world organised around attention, meaning, and recognition rather than production and consumption. The question of what gives life structure and purpose in a post-work society remained largely unresolved, revealing the limits of policy foresight when technological change disrupts fundamental assumptions about human motivation and social organisation.

Governance operates at a scale mismatch with technology. The CMEA scenario assumes coordinated global governance, but participants described that coordination as ‘reactive rather than proactive’ and ‘fragile, like Cold War-style nuclear cooperation’. Standards might be interoperable, but implementation would depend on national institutions with divergent capacities. Companies would become geopolitical actors, blurring the boundaries between market and state. The EU itself was discussed as both a potential model for supranational coordination and a politically constrained entity struggling to match the speed and scale of market-driven innovation coming from the US and China. The tension between global technological acceleration and slow demographic governance emerged as one of the most intractable problems in building a European ecosystem of excellence in AI. Coordination would prevent catastrophic outcomes but would not necessarily produce the legitimacy, capacity, or speed needed to manage large-scale labour displacement and maintain social cohesion.

The bundle 1 exercise revealed that the CMEA scenario's core challenge is not technological, but social, economic and institutional. Participants proposed many (technically and even politically) feasible policies across capital, R&D, data infrastructure,

talent, job quality, and inequality. Yet, the deeper difficulty lay in the architecture of decision-making itself: who has the authority to act at the necessary speed and scale. Existing democratic institutions might lack both the trust and capacity to implement sufficient and legitimate redistribution mechanisms. The key insight was that Europe's ability to build an ecosystem of excellence in AI under conditions of exponential change depends less on specific policy instruments, than on governance legitimacy, institutional reform, and the relationship between democratic processes and technological pace. These seem the conditions under which any foresight-driven policy must operate.

7. BUNDLED FUTURE 2: THE SLOWER, MISSION-ORIENTED SCENARIO

The FGLA bundle combines a macro scenario of *fragmented global governance* and *government-driven R&D* with a micro scenario of *logarithmic development* and *augmentation* of work. As with bundle 1, the backcasting exercise aimed to flesh out the world and identify policies to ‘make the most’ of this future – maximising opportunities and mitigating risks within the constraints posed by the four drivers. Participants explored how Europe could foster innovation despite slower technological growth and limited international coordination, steering AI to enhance human work, strengthen public sector capacity, and support socially beneficial outcomes. Before the discussion, the bundle was presented to participants as follows:

The FGLA bundle represents a world that is ‘slower’ and less technologically disruptive than in the CMEA bundle, but potentially more socially cohesive. It envisages a world in which governments lead AI R&D, directing innovation towards strategic public value missions such as health, education, and democratic resilience. AI capabilities advance steadily but without dramatic leaps, with a focus on augmenting human work rather than replacing it. In this context, global coordination is limited, as fragmented governance, competing national priorities, and geopolitical tensions produce divergent standards and slower cross-border collaboration. The absence of AGI lessens the need for global coordination. National governments or blocs tailor solutions to local contexts and values, and potentially compete among themselves (akin to the Cold War’s space race). Investment prioritises reskilling, work redesign, and domain-specific AI applications that enhance human expertise and service delivery, resulting in innovation that is slower but more inclusive and socially embedded, with labour remaining central to production and value creation.

7.1. WORKSHOP JOURNEY

Participants began by exploring and stabilising this baseline scenario, fleshing out in more detail the FGLA scenario. The discussion followed three main thematic channels: the reasons for the AI slowdown, the implications of government-led R&D, and the consequences for work and the labour market.

First, participants treated the reasons for slower capabilities development as preconditions for imagining the FGLA future. Understanding why AI growth would plateau was seen as essential before exploring the implications of other aspects of this scenario. In the absence of AGI, participants envisaged slower, yet steady, technological progress,

often described as ‘efficiency AI’ or ‘meagre AI’ – a technology that facilitates the performance of tasks and enhances efficiency without producing ‘big leaps’. Key – and at times concurrent – explanations for this development included:

- **Government-driven orientation.** AI development would be steered towards augmentation and public-value applications, slowing overall growth compared with a market-driven, efficiency-maximising path. Ethical oversight, sector-specific regulations (e.g. in health), and democratic checks would further constrain rapid development. Political cycles and public perceptions – where governments are seen as ‘wasting money’ – might also curtail R&D investment.
- **Technological and market constraints.** Participants discussed the possibility of an ‘AI bubble’, with private-sector R&D slowing in response to plateauing capabilities. Labour shortages and fractured global trade regimes could also reduce incentives for rapid AI expansion.
- **Resource and material constraints.** Assuming that the issues surrounding the energy- and resource-intensity of the technology are not solved, rising costs or increasing natural resource scarcity could create bottlenecks and decrease marginal returns, making further breakthroughs or even incremental improvements economically unviable.
- **Fractured global governance of AI.** Divergent values, political instability, state competition and geopolitical tensions could lead to splintered standards, duplication, smaller markets and less cross-border collaboration, slowing the rate of R&D and innovation.

Second, participants emphasised that government-driven R&D would steer AI development towards strategic societal missions, such as health or education, prioritising public-value objectives over purely commercial gains. This could channel AI towards socially beneficial applications, including ed-tech, health-tech, and governance tools, and contribute to long-term goals like disease prevention or environmental targets. However, some participants also cautioned that government-led AI could be applied to ‘law & order’ objectives – such as surveillance, militarisation, or border enforcement – and influence media ecosystems, raising ethical concerns. In this scenario, policymakers would prioritise human-AI collaboration and augmentation over automation, reflecting the plateauing of technological capabilities. Participants also debated whether a government-driven R&D approach might concentrate development in academic or public sector settings rather than commercial markets, and whether this would increase or reduce bureaucratic red tape.

Third, participants envisaged that employment would remain the dominant work situation, even though some workers would be ‘AI-native’ while others might struggle to adapt, increasing segmentation and inequality in the labour market. Another source of AI-driven inequality might be the emergence of a divide between ‘augmentable’ and ‘less augmentable’ workers and occupations. In particular, while service and office workers could benefit more from AI’s cognition-augmenting capabilities and productivity increases, augmentation for manual or physical jobs might be limited to safety-improving technologies such wearables or digital twins in logistics, manufacturing or construction.

Still, the white-collar sector was expected to shrink, especially in occupations using only ‘medium-level’ skills, with participants envisaging more work in the physical or social world. Yet, slower technological change would allow for smoother adaptation, and technology would never get good enough to supplant labour as a key source of value. At the same time, participants agreed that even incremental, logarithmic change would amount to something significant by 2045, raising questions about support for those unable to participate fully in the workforce, the design of early retirement schemes, and opportunities for new generations entering the labour market. Overall, innovation in this scenario is said to be producing less disruption.

As in the other group working on the CMEA scenario, this group encountered moments of analytical friction – this time around how to conceptualise *augmentation* versus *automation*. Participants found that even those forms of AI that augment human work could still lead to aggregate job losses through efficiency gains. The meaning of augmentation was seen as highly sector-dependent: in manufacturing it might primarily enhance safety and reduce physical risk, whereas in professional services it could accelerate output, raise quality, or expand client reach. These differences suggested that augmentation is a contextual process shaped by sectoral processes and task structures.

Participants further emphasised that effective augmentation cannot be designed in the abstract – it must be learned through practice, education, and up- and reskilling policies. They argued for sector-specific experimentation with AI applications, coupled with shared evaluation and data collection. The lessons learned from this experimentation could then inform broader mission-oriented R&D programmes to fuel innovation across the board in areas such as health or education.

After fleshing out the scenario, participants considered what an EU ecosystem for excellence in AI might look like in the FGLA world. The discussion focused on two fundamental pillars of such an ecosystem: AI development and AI governance.

Looking at *AI development*, the government-led nature of research would inevitably shape outcomes. Targeted funding mechanisms, akin to Horizon-style programmes, would channel resources primarily into academic research. Governments might still collaborate with selected companies, offering European high-quality private data to companies in exchange for specific objectives and fostering public-value AI solutions. The market landscape would likely be diverse, featuring many small companies operating in a competitive environment fuelled by government contracts and public procurement.

Participants also discussed the role of companies within this ecosystem: some could develop AI tools in-house, others could build services on top of foundational models, or rely on external AI-as-a-service providers, highlighting the growing importance of AI management within organisations. As centralised initiatives might emerge to coordinate compute resources, in a fragmented (and more geopolitically unstable) scenario, this could create security risks linked to cyberattacks.

The logarithmic pace of AI development would also influence resource and data use in two ways. First, energy consumption pressures would be moderated by efficiency improvements and diminishing demand, making it less of a problem in this scenario. Second, because AI systems would be less data-hungry, foundational models could operate without proprietary datasets, making large datasets much less strategically valuable.

On *governance*, an EU ecosystem of excellence would prioritise ‘trustworthy AI’, requiring high standards to satisfy public concern about sensitive data usage. This emphasis on responsible AI reflects both public demand and comparative advantage, enabling Europe to act independently in a world of fragmented global governance. It would offer a selling point to users, featuring AI that would be more reliable than that outside Europe. Policies would be shaped around the EU AI Act, which would continue to serve as a cornerstone of AI regulation – protecting democratic practices, supporting augmentation over automation, and maintaining Europe as a stable and safe hub for talent and innovation.

At the same time, with AI development slowing down, narratives attributing this outcome to ‘excessive regulation’ might become more prominent, putting the AI Act and related frameworks under internal pressure from lobbying and competing interests. This highlights the challenges of implementing robust governance even in a stable, government-driven innovation landscape.

7.2. POLICY ACTION AND UNDERLYING DYNAMICS

The discussion culminated in a backcasting exercise in which participants outlined possible policy avenues to maximise the opportunities of the FGLA scenario while mitigating the potential risks. It encompassed the constraints and implications associated with government-driven R&D, fragmented governance, logarithmic development of AI capabilities, and augmentation of work. The policy responses were clustered in five areas.

Building compute and data infrastructure. Participants highlighted the need for centrally governed compute capacity to support government-led research and niche AI companies. As noted above, energy efficiency improvements and reduced demand in a logarithmic development scenario would mitigate resource pressures, but maintaining sufficient infrastructure would still remain critical. Public-private collaboration would be central, facilitated by co-designed public policies that create the conditions for companies to take risks, innovate and scale – thus fostering skills development and deepening access to capital. Policymakers could explore targeted investments in hardware, energy supply, and secure facilities to sustain the EU's AI ecosystem and reduce vulnerability to fragmented governance and cyber risks.

Creating a diverse AI landscape. A core aim here would be to foster a rich landscape of smaller, specialised AI companies working in synergy with the public sector. Policy levers might include tax incentives, public infrastructure (including energy), mandatory interoperability requirements, and streamlined access to capital in order to stimulate niche innovation. Skills development, clear data valuation frameworks, and support for start-ups could enhance competitiveness and resilience. Legislation, including adaptations of the AI Act, could balance risk mitigation with responsibilities for governments to actively support company growth rather than solely restrict harmful uses.

Ensuring effective deployment in society. Effective deployment requires public trust and acceptance of AI. One way to enhance public trust would be to communicate more widely on high-profile public-value applications (e.g. health diagnostics) – clearly explaining the costs and benefits. Policies could include AI literacy initiatives for workers, SME owners and the public, as well as transparency requirements for AI models. By integrating AI into curricula and professional training, including Erasmus+ and other applied programmes, the EU could enhance uptake while building trust and reducing the perception of AI as a 'black box'.

Fostering adoption of AI to improve public services. Public services present a key domain for augmentation-focused AI. Participants discussed mandating AI applications in health, safety, and education, coupled with rigorous assessment of their effectiveness and their

overall quality. Subsidies or targeted support might be necessary to ensure availability and accessibility in public services, avoiding restrictions based on socio-economic status. Policies could also emphasise continual evaluation of AI outcomes, ensuring that deployment is inclusive and aligned with public-value objectives.

Addressing demographic shifts and labour market risks. In the FGLA scenario, augmentation-oriented AI would be coupled with an ageing European population, creating issues for labour market participation and social security funding. Early retirement schemes for workers who are close to retirement age and unable to acquire needed 'AI skills' might become a necessity, offset by voluntary extensions of working life and social security reforms. An overhaul of taxation systems could be envisaged, in which firms that make the most productivity gains from AI deployment are asked to contribute more. Special attention to AI developers and integrators as emerging occupational groups could guide workforce planning and reskilling initiatives.

In summary, the backcasting exercise highlighted how the EU would need policy avenues that span infrastructure, market development, societal deployment, public services, and social systems adaptation to navigate an FGLA future. These policies would reflect the interdependence of technological, institutional, and social factors. While challenges such as equitable access, public trust, and workforce transitions would remain, participants stressed that careful, coordinated interventions could enable Europe to leverage AI for public value, maintain democratic oversight, and foster a stable, inclusive AI ecosystem.

7.3. KEY TAKEAWAYS

The bundle 2 backcasting exercise explored the conceptual and institutional conditions under which Europe could sustain innovation and social cohesion in a slower, government-led AI future.

Mission-driven innovation alters the logic of progress. Participants viewed the logarithmic, government-oriented trajectory of AI as both stabilising and constraining. A slower technological advance would redirect R&D towards health, education, and democratic resilience, but raises questions about efficiency, incentives, and competition. Public missions could sustain innovation without acceleration, yet would depend on long-term funding cycles and institutional continuity, which political systems rarely maintain. The scenario would replace the fear of disruption with the risk of bureaucratic inertia.

Work would remain central but the augmentation-automation distinction was not seen as clear cut. The FGLA future would keep labour as a core source of value. The benefits of augmentation would nonetheless accrue unevenly: high-skilled professionals and AI-native workers would gain productivity and autonomy, while less augmentable roles would risk stagnation or exclusion, a proportion of them still at risk of automation. Mid-

skilled workers might still stand to lose out, with a hollowing out of the middle of the labour market. Participants imagined new divides between those who design and those who adapt to AI systems, demanding policies that integrate job redesign, lifelong learning, and recognition of non-automatable forms of work. Analytical friction over the augmentation-automation distinction remained, but participants recognised that favouring the former over the latter would involve sector-specific attention, experimentation and tailored policy intervention.

Fragmented governance would test state capacity and legitimacy. In this scenario, the challenge would not be runaway technology but limited global coordination. Government-driven R&D and national AI missions would require procurement competence, cybersecurity, and sustained public trust. Fragmentation would offer room for local experimentation but hinder scale and interoperability. The effectiveness of Europe's 'trustworthy AI' model would depend on its visible success in delivering public-value outcomes that people could recognise.

Responsible governance would be a source of excellence for Europe. In a slower, more fragmented global landscape, Europe's strength would lie in governance rather than breakthrough speed: steady mission-oriented funding, transparent evaluation, credible communication, and labour policies that share augmentation's gains. The FGLA scenario thus portrays a world where resilience, legitimacy, and social inclusion – not exponential growth – define excellence.

8. CONCLUSION

The foresight exercise mapped how Europe's AI ecosystem could evolve by 2045 and what policy choices would shape those paths. Desk research identified four drivers, which were used to build scenario skeletons. At the macro level, these were (i) global AI governance (fragmented vs coordinated) and (ii) the orientation of R&D on AI (market- vs government-driven). At the micro level, these were (iii) the trajectory of AI capability (exponential vs logarithmic) and (iv) the effect on work (automation vs augmentation).

On day 1 of the subsequent workshop, we split experts into a macro and a micro group to enrich the four scenarios on the macro and micro planes (8 scenarios in total, 4 micro and 4 macro). We ran a cross-group peer review to test internal logic. Then we linked the macro and micro drivers by voting on the most causally coherent pairings. On day 2 we reshuffled the groups and asked them to backcast from two selected, bundled futures to today's policy levers.

Three insights on the drivers stand out. *First*, global governance and R&D orientation set the boundary conditions for technological pace and workplace change. Who funds and directs AI – private capital under competition or public missions under constraint – shapes not only capability growth but how tools are deployed in firms and services. *Second*, speed alone is not destiny. Exponential trajectories amplify displacement risks but also expose gaps in redistribution and legitimacy; logarithmic trajectories reduce shock but raise risks of complacency, bureaucratic inertia, and uneven augmentation. *Third*, labour outcomes depend on organisation and policy, and not just on technical feasibility. Job design, bargaining power, evaluation capacity, and procurement choices determine whether AI substitutes or complements work.

The linking exercise produced two dominant causal chains: market-driven R&D aligned with exponential capability growth and automation; and government-driven R&D aligned with logarithmic growth and augmentation. To avoid caricature, the team inverted the remaining dimension, governance, to create two nuanced bundles for backcasting: CMEA (coordinated global governance + market-driven R&D + exponential development + automation) and FGLA (fragmented global governance + government-driven R&D + logarithmic development + augmentation).

Backcasting under CMEA highlighted structural tensions rather than quick fixes. Participants struggled with the economic logic of near-zero marginal costs alongside concentrated AI ownership. Other tensions included the legitimacy of large-scale redistribution (e.g. universal basic income) at the speed required, and the mismatch between exponential technological change and incremental democratic capacity. The policy catalogue was rich, but feasibility fell into three clusters. These were (i) easier near-

term steps (skill incentives and sectoral data standards); (ii) harder institutional shifts (worker voice and well-being metrics); and (iii) long-horizon reforms (capital-market integration and redesigned social contracts). The core challenge would be more institutional – authority, trust, and speed – than technical.

Backcasting under FGLA showed a different calculus. Mission-oriented, slower innovation could prioritise public value, steady adoption, and social cohesion. Yet it would depend on credible procurement, evaluation, cybersecurity, and communication to maintain legitimacy. Work would remain central but be stratified. More specifically, highly augmentable roles could gain the most, mid-skilled roles would risk stagnation, and manual work could accrue relative value. Fragmentation would limit cross-border scale but enable local experimentation. Resilience and sovereignty would replace global efficiency as system goals. Excellence would be institutional rather than technological, with reliable funding cycles, evidence-based scaling of what works, and labour policies that distribute the gains from augmentation.

Importantly, both groups experienced moments of analytical friction that revealed how even core concepts would become unstable under scrutiny. In the CMEA exercise, participants had difficulty with the ‘economic incoherence’ of a world in which artificial general intelligence drives production costs towards zero while profit and ownership structures persist. Redistribution proved equally problematic, raising doubts about legitimacy, capacity, and whether democratic systems could ever match the pace of exponential technological change.

In the FGLA exercise, friction centred on the blurred boundary between automation and augmentation. Even augmentative AI was expected to displace some jobs, with what counted as augmentation varying widely across sectors. Together, these tensions highlight that policy foresight must grapple not only with *what* to plan for, but also with what its categories truly *mean* in a world likely to differ profoundly from today’s. The pressure points themselves are where future governance challenges will surface.

Across both bundles the lesson is socio-economic and institutional, not technical. Europe’s success depends on governance capacity and legitimacy to steer AI under very different macro conditions. In the CMEA scenario, coordinated rules could not on their own offset rapid displacement or resolve redistribution at the speed of exponential change. Politically feasible measures cluster in skills incentives and standards, while deeper reforms to markets and social protection face time and trust constraints.

In FGLA, slower, mission-oriented R&D could deliver public value only if governments are able to sustain credible procurement, protect security in a fragmented global landscape,

and continually evaluate whether AI deployments improve social outcomes for people. The actionable priority is therefore to match instruments to context. This means pairing visible, public-facing use cases with evaluation, using procurement and funding cycles to scale what works, and linking productivity gains to inclusion through representation, reskilling, and fairer distribution.



CEPS
PLACE DU CONGRES 1
B-1000 BRUSSELS
