

European Chips Act : "Les 27 États membres poursuivent des priorités divergentes."

Avertissement : Souveraine Tech revendique par vocation une approche transpartisane. Seule nous oblige la défense des intérêts supérieurs de notre pays. Nous proposons ainsi un lieu de « disputatio » ouvert aux grandes figures actives de tous horizons. La parole y est naturellement libre et n'engage que ceux qui la prennent ici. Cependant, nous sommes bien conscients des enjeux en présence, et peu dupes des habiles moyens d'influence plus ou moins visibles parfois mis en œuvre, et dont tout un chacun peut faire l'objet, ici comme ailleurs. Nous tenons la capacité de discernement de notre lectorat en une telle estime que nous le laissons seul juge de l'adéquation entre le dire et l'agir de nos invités.

[Thibault Laurentjoye](#) est « Associate Professor » à l'[Aalborg University Business School](#) ☐☐

Cet entretien a été publié le vendredi 23 janvier 2026.

1a/ Pourquoi le passage des architectures planaires à FinFET puis GAA a-t-il transformé les semi-conducteurs en un enjeu de puissance étatique, et pas seulement en une question d'ingénierie ?

Jusqu'aux années 2000, améliorer les processeurs consistait simplement à réduire la taille de transistors planaires, c'est-à-dire des architectures en deux dimensions où le passage du courant à travers un canal est déterminé par la

tension dans une grille. Cependant, au fur et à mesure la miniaturisation a généré des effets électrostatiques perturbateurs, imposant un véritable changement de paradigme technologique : le passage au FinFET (canal en forme d'aileron – « fin » en anglais), une architecture tridimensionnelle augmentant la surface de contact entre le canal et la grille.

Cette transition a permis une densification spectaculaire : on atteint aujourd'hui entre 250 et 300 millions de transistors par millimètre carré sur les puces dites à 3 nanomètres. Cette densité a rendu physiquement possible l'intelligence artificielle moderne, qu'il s'agisse des modèles de langage large comme GPT, Claude, DeepSeek ou Gemini, ou des modèles de diffusion pour la vidéo, l'image et le son (Sora, Suno, Nano Banana).

Ces applications touchent désormais des domaines régaliens : défense (systèmes d'armes autonomes, cybersécurité), renseignement (analyse de données massives), santé publique (diagnostic médical assisté par IA), ou encore souveraineté numérique (clouds de confiance, infrastructures critiques). Pour un État comme la France, ne pas maîtriser cette technologie, c'est dépendre d'acteurs étrangers pour des fonctions essentielles à sa sécurité nationale : d'où l'enjeu de puissance étatique.

1b/ Que révèle cette complexification sur les barrières à l'entrée et la dépendance vis-à-vis de quelques acteurs clés ?

Cette complexification technologique a créé une concentration industrielle inédite. Intel, pourtant inventeur du FinFET, a échoué à le rentabiliser par manque de cohérence stratégique interne. Seuls TSMC à Taïwan et Samsung en Corée du Sud ont réussi cette transition critique et dominant aujourd'hui la production avancée de semi-conducteurs. Le secteur privé refuse d'assumer seul les risques colossaux liés à ces investissements : une fab moderne coûte environ 20 milliards

de dollars, pour prendre l'exemple de celle construite en Arizona à la suite du CHIPS Act en 2022. Cette réalité crée des barrières à l'entrée formidables. Le coût unitaire d'une machine de lithographie EUV se compte en centaines de millions de dollars, allant de 100 à 400 millions selon la génération, la dernière en date – dite à « haute ouverture numérique », qui va être nécessaire pour la production de puces 2 nm – valant plus que le PIB de certains petits pays. Ces montants considérables expliquent l'envolée du coût des fabs et révèlent une dépendance géopolitique extrême envers deux acteurs asiatiques pour toute l'économie numérique mondiale. (Il faut rappeler que la Corée du Sud est également le leader mondiale en matière de mémoire vive, à travers Samsung et SK Hynix.)

2a/ En quoi la lithographie EUV constitue-t-elle un « verrou stratégique » dans l'accès aux nœuds avancés (<7 nm), et pourquoi l'Europe y occupe-t-elle une position à la fois centrale et fragile ?

Pour produire des puces en dessous de 7 nanomètres avec un bon rendement, c'est-à-dire sans que la majorité de la production parte au rebut, la lithographie à ultraviolets extrêmes (EUV) est devenue incontournable. Or, une seule firme au monde fabrique ces machines : le néerlandais ASML. Auparavant, c'est-à-dire jusqu'à 10nm, on pouvait utiliser la lithographie à ultraviolets profonds, technologie pour laquelle il existait davantage de concurrence – on trouvait ainsi Canon ou Nikon.

Le monopole technologique d'ASML sur la lithographie EUV doit être remplacé au sein d'un écosystème incluant notamment l'allemand Zeiss qui fabrique les miroirs les plus parfaits du monde. L'EUV est un goulot d'étranglement absolu : sans accès à ASML, impossible de produire des semi-conducteurs avancés. L'Europe occupe ainsi une position centrale puisqu'ASML est européenne, mais fragile car cette entreprise dépend également de la propriété intellectuelle étatsunienne et compte Intel parmi ses actionnaires. De façon générale, on peut dire que

l'Europe n'utilise pas son pouvoir de négociation et reste dans une approche naïve de l'intégration dans l'économie mondiale.

2b/ Quels enseignements tirer pour une politique industrielle européenne réaliste ?

Le grand paradoxe européen, qui est au centre du diagnostic que je pose dans ma note, réside dans le fait que l'Europe possède d'une part, une excellence remarquable en recherche et développement, avec des acteurs phénoménaux comme le CEA-Leti en France, IMEC en Belgique et Fraunhofer en Allemagne, qui sont des pôles d'excellence mondiaux ; d'autre part, à l'autre extrémité de la chaîne de valeur, l'Europe dispose d'un marché mature de centaines de millions de consommateurs de semi-conducteurs avancés dans les smartphones, les ordinateurs et de plus en plus dans l'IA ; mais entre les deux, on n'a rien.

Sur la partie avancée des semi-conducteurs, disons en dessous de 12 nm, on ne produit à peu près rien, à l'exception d'un site Intel en Irlande qui exporte toute sa production vers les États-Unis. Cette absence du milieu de la chaîne de valeur des semiconducteurs avancés laisse l'Europe totalement exposée en cas de perturbation. L'enseignement principal est que l'Europe n'utilise absolument pas son pouvoir de négociation et reste prisonnière d'une naïveté ricardienne fondée sur les avantages comparatifs, ignorant les contraintes géopolitiques. Une politique réaliste devrait viser une autosuffisance de 10 à 50% selon les créneaux, notamment sur le segment 5 nanomètres pour l'inférence IA, en passant par des banques promotionnelles comme la BEI, BPI France ou KfW plutôt que par des subventions directes.

3a/ Pourquoi le développement rapide de l'IA renforce-t-il une dépendance extrême aux nœuds les plus avancés, et quels risques cela crée-t-il pour les économies qui n'y ont pas accès direct ?

L'intelligence artificielle comporte deux phases distinctes

qu'il est crucial de distinguer : l'apprentissage et l'inférence. L'apprentissage – qui est la phase d'entraînement des modèles – nécessite les technologies les plus à la pointe, typiquement 3 ou 4 nm dans le contexte actuel. L'inférence – l'utilisation des modèles par les utilisateurs comme vous et moi – peut fonctionner sur des structures moins avancées, typiquement sur du 5 à 7 nanomètres – voire 10-12 nm dans le cas de modèles plus petits, ayant vocation à tourner sur des configurations locales.

À l'exception du dernier Gemini 3 qui a sans doute été entraîné sur du 3 nm, presque tous les modèles récents que tout le monde utilise, qu'il s'agisse de GPT, Claude, DeepSeek ou d'autres, ont été entraînés sur du 4 nm, car leur matériel était fourni par Nvidia, en particulier la série d'accélérateurs Blackwell. Nvidia était sur du 4 nanomètres et va passer avec Rubin sur du 3 nanomètres. Cependant, que l'on parle de Google ou Nvidia, tout le monde s'approvisionne chez TSMC, qui jouit d'un quasi-monopole de fait sur la production de semiconducteurs de 4 et 3 nm.

Aujourd'hui, il n'y a que 4 pays (ou zones) qui produisent des puces de taille inférieure ou égale à 7nm : Taiwan avec TSMC, la Corée du Sud avec Samsung, les Etats-Unis avec Intel (on peut y ajouter le site irlandais d'Intel qui est de facto un satellite US) et la Chine avec SMIC – la différence majeure étant que SMIC est le seul de ces acteurs à ne pas avoir accès aux machines EUV produites par ASML. On l'a vu récemment avec DeepSeek pour leur modèle R2, il est à peu près impossible pour eux de se passer de matériel Nvidia pour la phase d'apprentissage (Google n'étant pas une option). Par contre, une fois que le modèle est entraîné, les puces produites par SMIC montées sur les accélérateurs Ascend de Huawei font tout à fait l'affaire pour l'inférence.

3b/ Peut-on imaginer des trajectoires d'IA moins

dépendantes du very-advanced silicon ?

Oui, en distinguant clairement les besoins pour l'apprentissage et pour l'inférence. Mais l'Europe ne part pas de rien : elle possède Mistral, le seul modèle de langage large non américain et non chinois dans le top 10 mondial. C'est un atout considérable. Cependant, Mistral dépend largement d'infrastructures américaines pour l'entraînement de ses modèles, ce qui crée une vulnérabilité stratégique.

La question est de savoir par où commencer. Produire d'emblée du 3 et 4 nanomètres serait probablement trop ambitieux quand on part de très peu. L'Europe pourrait commencer par du 5 nanomètres, car les besoins en matière d'inférence sont clairs et en croissance exponentielle. Cela permettrait à Mistral de déployer ses modèles sur une infrastructure au moins partiellement européenne.

À plus long terme, un arbitrage devra se poser entre usage croissant de ressources matérielles et énergétiques d'une part, et raffinement et optimisation des modèles existants d'autre part. Aujourd'hui, la Silicon Valley est un peu dans une course qui ressemble à une fuite en avant. L'Europe pourrait adopter une trajectoire différente, plus sobre et stratégiquement ciblée.

4a/ Comment la spécialisation croissante du hardware (GPU, TPU, accélérateurs dédiés) modifie-t-elle la chaîne de valeur mondiale et la capacité des États à « rattraper » technologiquement leur retard ?

Pour faire simple, trois types de processeurs coexistent désormais dans l'écosystème de l'IA. Les CPU, processeurs centraux des ordinateurs, sont les plus puissants en termes de puissance brute mais ne peuvent gérer qu'un nombre réduit de processus à la fois. Ils sont hautement spécialisés avec une puissance très concentrée.

Les GPU, processeurs graphiques, sont plus adaptés à effectuer

des calculs plus nombreux. Ils sont moins forts en puissance brute que les CPU mais peuvent réaliser davantage de calculs parallèles. Vu que la plupart de l'intelligence artificielle repose sur des calculs matriciels, pouvoir casser à un moment donné le calcul principal en différents sous-calculs afin de réaliser les opérations matricielles est quelque chose que le GPU fait mieux que le CPU.

Cependant, un nouveau joueur est arrivé : le TPU de Google, où T signifie tenseur. Il s'agit d'un architecteur de processeur ultra-spécialisée pour le calcul matriciel et encore plus efficace que le GPU, y compris sur le plan énergétique. Google l'a développé en interne pour l'apprentissage et l'inférence de ses modèles Gemini. Apparemment, Anthropic, l'entreprise qui développe Claude, a commencé à nouer un contact avec Google pour utiliser aussi des TPU.

4b/ La spécialisation est-elle une opportunité ou un piège pour l'Europe ?

La spécialisation du hardware représente à la fois une opportunité et un piège redoutable pour l'Europe. Le piège réside dans le renforcement de la concentration industrielle. TSMC contrôle 95% du marché 3 nanomètres après l'échec de Samsung dans sa transition vers le Gate-All-Around. Samsung avait tenté de passer directement au GAA mais a rencontré de mauvais rendements, rendant la production non viable commercialement. TSMC, en restant prudemment sur FinFET pour le 3 nanomètres, a complètement empoché la mise. Ils ne passeront au GAA qu'avec le 2 nanomètres, dont ils débutent tout juste la production à grande échelle. Cette concentration signifie que toute l'infrastructure d'IA mondiale, que ce soient les GPU Nvidia ou les TPU Google, dépend d'un seul acteur taïwanais. Pour l'Europe sans capacités de production avancée, c'est une dépendance totale.

Cependant, l'opportunité existe : la spécialisation ouvre des

niches stratégiques. Les TPU de Google montrent qu'on peut innover en hardware spécialisé sans nécessairement maîtriser toute la chaîne. L'Europe pourrait développer des accélérateurs dédiés pour des applications spécifiques comme la santé, le climat ou la défense, sur des nœuds légèrement moins avancés, entre 5 et 10 nanomètres, où la dépendance est moindre. Viser l'autonomie complète serait ruineux et probablement vain, tandis que miser sur une puissance sélective et stratégique pourrait être viable.

5a/ Pourquoi l'explosion des coûts fixes des fabs avancées conduit-elle à une concentration industrielle extrême, et quelles conséquences cela a-t-il sur la souveraineté technologique nationale ?

Une fab moderne pour semi-conducteurs avancés coûte environ 20 milliards de dollars, comme celle construite en Arizona. Cet investissement colossal s'explique notamment par l'achat obligatoire de plusieurs machines EUV à ASML, dont le prix unitaire se compte en centaines de millions de dollars. Dès qu'on veut produire des semi-conducteurs avancés, il faut acheter quelques machines à ASML et rien que cela fait passer facilement au-delà du milliard. Face à de tels montants et à une certitude insuffisante sur le retour sur investissement à moyen terme, le secteur privé refuse d'investir seul. Les acteurs privés ont tendance à être réticents face à certaines formes de risques, et investir 20 milliards avec une incertitude à moyen terme insuffisante est source d'un refus d'obstacle.

Seuls quelques acteurs bénéficiant d'un soutien étatique massif ou d'une position dominante peuvent franchir cette barrière, créant une concentration industrielle sans précédent et des conséquences majeures sur la souveraineté technologique nationale. La conséquence en matière de souveraineté technologique pour les pays ne disposant pas de sites de production, est que soit ils vont être dépendants d'autres pays pour accéder à des certaines technologies critiques, soit

ne pas y avoir accès.

5b/ Jusqu'où un État peut-il (ou doit-il) subventionner cette industrie ?

Les subventions au sens large sont une nécessité pour les industries naissantes ou renaissantes. List parlait de protectionnisme éducateur. Le secteur des semiconducteurs ne fait pas exception à cette règle : TSMC comme Samsung ont bénéficié des concours financiers de leur gouvernement respectif.

Cependant, dans un contexte de contraintes budgétaires, comme celles des pays de la zone euro, on n'est pas obligé de passer par de pures subventions, où l'État donnerait directement de l'argent. On pourrait passer par un montage avec des banques dites promotionnelles, les banques publiques comme la Banque européenne d'investissement (BEI), BPI France en France ou KfW en Allemagne. De bonnes conditions de financement dans le contexte d'un partenariat public-privé stratégique constituent un équivalent de subventions qui apparaît moins problématique pour des finances publiques contraintes.

Jusqu'où faut-il subventionner ? Idéalement, jusqu'à atteindre une forme de taille critique, c'est-à-dire une stabilisation des coûts moyens de production à un niveau suffisamment compétitif. Il faut subventionner pour montrer aux agents privés que l'État est là, qu'il y a une forme de garantie qui est donnée, doublée d'un signal d'engagement des pouvoirs publics en faveur du développement d'une certaine industrie. Tant qu'on ne fera pas cela, je suis assez pessimiste sur la capacité européenne en général, et française en particulier, à rebondir en matière industrielle. Le montant du CHIPS Act américain versus européen illustre cette différence d'ambition et explique l'écart de résultats.

6/ Comment les pénuries de semi-conducteurs se

transmettent-elles à l'ensemble de l'économie (automobile, énergie, biens durables), et que nous apprend l'épisode post-COVID sur la vulnérabilité macroéconomique européenne ?

L'industrie automobile, touchée durant la période du COVID, illustre parfaitement cette transmission. Elle utilise de nombreuses puces de différentes tailles, et si une seule taille manque, dont le prix n'est que de quelques euros seulement, c'est un véhicule entier qui coûte des dizaines de milliers d'euros qui ne peut pas être produit. Ces facteurs sont complémentaires et non substituables, ce qui fait d'ailleurs que les modèles d'économie mainstream ont eu du mal à vraiment appréhender la réalité des chaînes de valeur et des perturbations qu'on connaissait.

L'épisode post-COVID a révélé notre vulnérabilité à travers les ruptures créées par les confinements non synchronisés dans les différentes zones du monde, avec des périodes d'ouverture en Europe et de fermeture en Asie. Les producteurs automobiles avaient réduit leurs commandes de semi-conducteurs pendant la pandémie, vu que leur propre demande avait baissé. Les places libérées dans la production de puces sont allées vers les compagnies dont la demande bondissait, notamment toutes celles ayant des activités en ligne, créant une augmentation énorme de la demande en composant liés aux serveurs. Lorsque les firmes automobiles se sont réveillées, les places de production étaient saturées et elles ont dû attendre beaucoup de temps, bloquant toute l'industrie automobile pendant des mois.

7a/ En quoi les pénuries de puces ont-elles contribué à l'inflation récente, et pourquoi cette dimension « technologique » est souvent absente des débats économiques classiques ?

Les goulots d'étranglement créés par les pénuries de semi-conducteurs ont créé deux mécanismes inflationnistes. D'abord, dans certains cas désespérés, on a été obligé de trouver des

alternatives, mais la recherche d'alternatives avait un coût. Dans d'autres cas, cela a simplement créé des pénuries et donc des effets d'enchère sur la rareté. Il faut se souvenir que le premier poste aux États-Unis qui a vu une augmentation significative des prix fut les voitures d'occasion. Lorsque l'économie étatsunienne a redémarré en 2021, le fait qu'il n'y avait pas de voitures neuves a fait bondir le prix des voitures d'occasion. Or, la pénurie de voitures neuves était grandement liée aux semi-conducteurs. Cette dimension technologique de l'inflation, où des composants à quelques euros bloquent des biens finaux à dizaines de milliers d'euros, reste encore largement absente des débats économiques classiques centrés sur la demande agrégée et la politique monétaire.

7b/ Comment mieux intégrer les contraintes matérielles dans les politiques économiques ?

Les modèles de type entrées-sorties (input-output) avec des intrants physiques, comme par exemple les analyses en termes de cycle de vie des produits et des chaînes de valeur, sont sans doute les approches les plus intéressantes qu'on a actuellement. Elles mélangent approches économiques et physiques dans une perspective holistique, et c'est sans doute quelque chose dont on a vraiment besoin pour retrouver cette articulation entre planification et dynamique de marché. Cartographier les chaînes de valeur pour s'assurer que les intrants nécessaires sont localisables et sourçables devrait être un prérequis à toute politique industrielle.

8a/ Les semi-conducteurs peuvent-ils devenir un goulot d'étranglement majeur pour la diffusion de l'IA et la transition énergétique (électrification, réseaux, stockage) ?

Absolument. Plus l'IA deviendra centrale dans l'organisation administrative et productive, plus la dépendance aux semi-conducteurs avancés s'accroîtra. Le fait qu'il n'y ait pas de chaîne d'approvisionnement, de sites de production industriels

sur des tailles inférieures à 12 nanomètres en Europe, est une faiblesse importante.

Pour la transition énergétique à proprement parler, les composants de base de l'électrification utilisent des puces entre 20 et 60 nanomètres, ce qui est moins critique. De ce point de vue, on est dans une passe un peu meilleure. Cependant, les systèmes intelligents de gestion qui maximisent l'efficacité de ces infrastructures – réseaux électriques intelligents, optimisation algorithmique de la distribution énergétique, systèmes avancés de gestion des batteries – bénéficient de puces plus performantes. À mesure que l'IA s'intègre dans la transition énergétique, les besoins en semi-conducteurs avancés augmenteront là aussi.

En cas de perturbation d'origine géopolitique, sanitaire ou autre, l'Europe serait totalement paralysée, une vulnérabilité stratégique majeure face aux ambitions climatiques et numériques affichées.

8b/ Quels arbitrages stratégiques cela impose-t-il aux décideurs publics ?

Trois arbitrages stratégiques auraient du sens pour les décideurs publics. Premièrement, concernant la transition énergétique versus l'IA : pour l'électrification et les réseaux intelligents, les puces nécessaires sont du 20 à 60 nanomètres. Cependant, l'optimisation algorithmique via l'IA nécessitera du silicium plus avancé, et ce serait une bonne idée de produire nous-mêmes les puces qui vont dans la gestion de certains réseaux d'électricité. Deuxièmement, l'arbitrage entre apprentissage et inférence : se concentrer sur l'inférence qui nécessite du 5 à 10 nanomètres plutôt que sur l'apprentissage qui nécessite du 3 à 4 nanomètres. Troisièmement, l'arbitrage sur l'autosuffisance : viser entre 10 et 50% d'autosuffisance selon les créneaux, notamment pour Mistral. Il serait important que l'inférence du champion européen tourne au moins partiellement sur une infrastructure

européenne, et notamment française. Garantir qu'une partie des composants de la chaîne de valeur sont produits en Europe serait déjà une amélioration considérable.

9a/ Le retard technologique chinois dans les nœuds avancés est-il principalement conjoncturel (lié aux sanctions) ou structurel (lié aux limites du modèle industriel et scientifique) ?

Le retard chinois en matière de semi-conducteurs est principalement conjoncturel et lié aux sanctions bloquant l'accès aux technologies clés. Ce retard est d'autant plus notable qu'il constitue l'exception qui confirme la règle : aujourd'hui, la Chine est en train de devenir le leader mondial industriel incontestable. Sur des technologies liées aux énergies renouvelables, le solaire, l'éolien, les batteries, la Chine a une avance considérable. Elle a au moins une présence en matière de fabrication, et de plus en plus la propriété intellectuelle. La Chine est passée en à peine deux décennies d'un suiveur à un meneur sur le plan technologique. Les semi-conducteurs constituent l'une des rares exceptions à cette règle, et en conséquence l'IA également, même si avec DeepSeek, Qwen ou Kimi, la Chine demeure le principal rival des États-Unis.

Le retard chinois en matière de semiconducteurs de pointe est dû à l'impossibilité d'avoir accès aux machines de lithographie EUV produites par ASML. L'entreprise chinoise SMIC, qui est un peu le TSMC local, a du mal à passer en dessous de 7 nanomètres. Ils arrivent déjà à faire du 7 nanomètres sans ultraviolet extrême, ce qui en soi est un exploit que seul TSMC avait réussi à réaliser il y a plusieurs années, mais ils ne peuvent pas passer à 5 ou 6 nanomètres sans EUV, c'est presque physiquement impossible. Ils essaient de compenser, à travers des passages multiples en opérant des micro-déplacements entre deux passages, mais les rendements sont mauvais et donc ce n'est pas viable économiquement.

Cependant, derrière SMIC il y a Huawei, une entreprise extrêmement sérieuse. Huawei met les moyens pour développer leur propre technologie d'ultraviolets extrêmes. Cette technologie avait mis plusieurs décennies à être développée – Chris Miller évoque entre 25 et 30 ans de recherche dans son ouvrage *Chip Wars*. Pour accélérer le processus, Huawei cherche à recruter d'anciens employés d'ASML ou des personnes qui ont travaillé sur les machines EUV, pour tenter de réunir suffisamment d'intelligence et de compréhension du sujet afin de concevoir leur propre technologie.

De mon point de vue, la question n'est pas de savoir s'ils y arriveront, mais quand ils y arriveront. On peut considérer que ce rattrapage leur prendra probablement entre 5 et 10 ans, sur la base d'un développement interne relativement autonome.

9b/ Que cela implique-t-il pour la durée et l'intensité de la rivalité États-Unis / Chine ?

Durant la prochaine décennie, les États-Unis conserveront vraisemblablement leur avance technologique en matière de semiconducteurs et d'IA. Cependant, la Chine déploie déjà des trésors d'inventivité pour optimiser les algorithmes et les méthodes de fonctionnement – on l'a encore vu avec le papier récent de DeepSeek sur les hyper-connexions contraintes.

Cette situation met en lumière le rôle incroyablement stratégique de Taïwan et montre à quel point le monde est dépendant d'une île dont le statut géopolitique est contesté. Même si les experts s'accordent à dire qu'une invasion militaire de Taïwan est extrêmement peu probable pour des raisons logistiques, ainsi qu'en raison du coût humain et financier potentiel pour la Chine, en revanche un blocus autour de Taïwan serait envisageable et pourrait perturber les chaînes de valeur de façon considérable. Dans ce scénario un peu catastrophe pour l'Europe, avec des États-Unis relativement peu coopératifs comme c'est le cas actuellement, on serait dans une situation potentiellement assez difficile,

du fait de l'absence de production domestique minimale pour alimenter nos usages les plus stratégiques.

10a/ Pourquoi l'Europe, malgré une excellence en R&D (France, Belgique, Pays-Bas), reste-t-elle structurellement faible sur la production industrielle avancée ?

Le principal problème de la politique industrielle européenne, c'est qu'elle n'existe pas. C'est un peu le péché originel de l'Union européenne d'avoir décrédibilisé l'action et la vision de l'État. Cela se voit dans les règles budgétaires qu'on a, qui sont complètement ineptes et qui ont limité les marges de manœuvre des États, faisant des dépenses d'investissement public la variable d'ajustement par défaut – au détriment de l'avenir du continent et des pays qui le composent.

En Europe, on peut dire qu'il n'y a pas eu de véritable création autonome de nouveaux pôles d'excellence industrielle au cours des dernières décennies. Ce qui existe en matière de politique industrielle existait déjà il y a plusieurs décennies. Il n'y a pas eu d'amélioration, il n'y a pas eu de création de nouveaux pôles d'excellence industrielle véritable en Europe depuis l'introduction de l'euro, par exemple. Au mieux, on a amélioré l'existant.

En conséquence, lorsque l'on compare aux États-Unis, mais surtout à la Chine où la notion d'écosystème industriel est absolument au centre du modèle de développement industriel, l'Europe est vraiment à la ramasse. Le Chips Act européen a des montants qui n'ont absolument rien à voir avec le CHIPS Act américain. Cette faiblesse structurelle découle d'une absence de vision stratégique et d'une naïveté géopolitique persistante face aux réalités du vingt-et-unième siècle, restant prisonnière d'une vision ricardienne des avantages comparatifs et d'une espèce de naïveté du doux commerce de Montesquieu poussé à son extrême, au détriment de toute prise en compte des contraintes géopolitiques.

10b/ Quelles sont les limites réelles du Chips Act européen face au CHIPS Act américain ?

Le European Chips Act souffre de trois faiblesses structurelles qui compromettent son ambition affichée de porter l'Europe à 20% de la production mondiale d'ici 2030.

Premièrement, l'écart d'échelle avec les États-Unis est massif. Le chiffre annoncé de 43 milliards d'euros est trompeur : il agrège seulement 3,3 milliards de budget européen direct, environ 12 milliards de subventions nationales, et le reste en investissements privés associés. En proportion du PIB, un effort équivalent au CHIPS Act américain nécessiterait 48 milliards d'euros de fonds publics – soit plus du triple de ce qui a été effectivement mobilisé.

Deuxièmement, le ciblage technologique est déséquilibré. Les projets approuvés à ce jour (ESMC Dresden, STMicro-GlobalFoundries Crolles, Infineon Dresden) concernent quasi exclusivement des nœuds matures (12nm et au-dessus) et essentiellement destinés à l'automobile. C'est bien, mais cela laisse complètement de côté les nœuds plus avancés, particulièrement inférieurs 7nm qui alimentent l'IA et le calcul haute performance. Le retrait d'Intel de son projet 2nm à Magdeburg laisse l'Europe sans aucune initiative crédible sur ce segment. Le projet d'Intel était mal ficelé et paradoxalement trop ambitieux, mais son arrêt laisse un vide béant dans la stratégie européenne en matière de semiconducteurs.

Troisièmement, la gouvernance reste fragmentée. Les 27 États membres poursuivent des priorités divergentes, sans mécanisme de coordination comparable au Department of Commerce américain. Les subventions nationales sont soumises aux règles d'aides d'État, créant des délais et des incertitudes juridiques dont on se serait bien passés.

En somme, le Chips Act européen mobilise des ressources insuffisantes, omettant des segments stratégiques, avec une

coordination défaillante – perpétuant plutôt qu'il ne résout le paradoxe de la chaîne d'approvisionnement européenne.

10c/ Faut-il viser l'autonomie complète ou une puissance sélective et stratégique ?

L'autonomie complète est hors de portée. Même Samsung a eu du mal à produire du 3 nanomètres à grande échelle de façon viable. Il faut être réaliste. On ne va pas internaliser l'intégralité de la chaîne de valeur. Déjà dans l'absolu, cela n'aurait pas forcément un grand sens, et en tout cas on n'arrivera pas à le faire rapidement. Cependant, essayer d'avoir un minimum de présence dans les différentes étapes serait une forme de bon sens.

Pour des raisons stratégiques, TSMC évite de laisser partir en temps réel les technologies les plus avancées. Par exemple, ce qu'ils ont fait aux États-Unis, ils ne diffusent sur des sites de production extérieurs que la seconde technologie la plus avancée. Ils envisagent même de ne déléguer que la troisième technologie la plus avancée à partir de l'an prochain. On pourrait donc négocier avec eux, pour produire du 5 nanomètres, que l'on ne produit pas du tout actuellement, alors qu'on en aurait besoin pour l'inférence en matière d'IA, dont les besoins croissent de façon exponentielle.

La stratégie doit être une puissance sélective avec une autosuffisance à présence variable selon les créneaux. On a la chance d'avoir Mistral. Il serait important que, pour Mistral au moins, on puisse faire tourner l'inférence sur des composants dont une partie de la chaîne de valeur est produite en Europe et même potentiellement en France. Cela serait une amélioration significative. L'objectif n'est pas de tout faire, mais de ne pas être à zéro pour éviter la paralysie totale lors de perturbations géopolitiques ou sanitaires.