



FONDATION
FRANCE-ASIE



FRANCE INDIA
FOUNDATION

Indo-French Perspectives on Artificial Intelligence

White Paper
First Recommendations

France India AI Initiative

February 2026

The France India AI Initiative white paper
is endorsed by Anne Bouverot

France's Special Envoy for AI
Chair of the Board of Ecole Normale Supérieure



**FONDATION
FRANCE-ASIE**

The Fondation France-Asie
is an independent foundation dedicated to
relations between France and major Asian countries.

Launched in 2012, the Fondation France-Asie encourages
dialogue and the development of new partnerships
between France and Asia in the service of shared values.

The Young Leaders Programs, at the heart of its mission, are
organized with the various country chapters with
India, China and Japan.

Nicolas Macquin, Co-founder & President

Arnaud Ventura, Co-founder

Thomas Mulhaupt, Managing Director



**FRANCE INDIA
FOUNDATION**

The France India Foundation
is an independent organisation dedicated to strengthening
Franco-Indian collaborations across diverse sectors through
the promotion of meaningful dialogue and facilitating
interpersonal connections and engagements.

As an integral part of the France Asia Foundation network, the
France India Fondation concentrates its efforts on the Indian
region. Working in tandem with the
Fondation France-Asie leadership,
it actively executes its vision in India to drive impactful
initiatives and forge lasting connections.

Sumeet Anand, Co-founder & President
Chiki Sarkar, Co-founder
Payal S. Kanwar, Representative General Director

EXECUTIVE SUMMARY	7
KEY RECOMMENDATIONS	10
CONTEXT & INTRODUCTION	12
MILESTONES OF THE INITIATIVE	20
CONTRIBUTORS OF THE INITIATIVE	22
AI & HEALTHCARE	26
1. Strategic Review and Priorities	27
2. Scope	28
3. Observations	29
4. Practical illustrations & Use cases	30
5. Recommendations	39
AI & AUTOMOTIVE	46
1. Strategic Review and Priorities	46
2. Scope	47
3. Observations	49
4. Use cases	51
5. Recommendations	53
AI & REGULATION	62
1. Strategic Review and Priorities	62
2. Recommendations	81
CONCLUSION	85
CORPORATE PARTNERS	86
EXECUTIVE TEAM	88
CONTACTS	88

EXECUTIVE SUMMARY

This white paper is a first joint contribution of the France–India AI Initiative. It should be read as an initial working document designed to structure dialogue, identify shared priorities, and formulate early recommendations to support sustained and results-oriented cooperation between France and India in artificial intelligence.

Artificial intelligence is reshaping scientific research, industrial systems, and public services while raising interconnected challenges related to data governance, validation, trust, skills, security, infrastructure, and sustainability. In this evolving global landscape, Franco-Indian cooperation offers distinctive potential. France contributes recognised strengths in research, regulation, and high-value industrial sectors, while India brings scale, digital public infrastructure, engineering talent, and deployment capacity. These complementarities create a practical foundation for cooperation capable of linking innovation with implementation and trust frameworks with measurable impact.

This first document proposes an initial cooperation framework structured around three priority domains where joint action can generate both short-term and structural impact: AI & Healthcare, AI & Automotive, and AI & Regulation. These workstreams combine high-impact application sectors with a transversal governance dimension and are intended as practical entry points for progressive collaboration.

AI & Healthcare

The paper argues that artificial intelligence in healthcare should be approached as public-interest infrastructure rather than as a purely technological competition. Ensuring trust, clinical safety, sustainability, and equitable outcomes requires robust governance, validated deployment pathways, and strong professional adoption. Brain health is identified as a flagship domain for initial Indo-French cooperation given its scientific complexity and governance sensitivity. It provides a rigorous testbed for responsible medical AI deployment before expansion to broader global health applications. France and India are well positioned to implement reliable and trustworthy medical AI use cases within the next 12–36 months by leveraging complementary health data ecosystems, institutional frameworks, and research capacities.

Key structural barriers remain: fragmented governance and validation pathways, constraints on cross-border data collaboration, gaps between research and clinical deployment, limited training and trust among health professionals, risks of bias at population scale, and insufficiently structured long-term cooperation frameworks.

The paper therefore recommends coordinated action to establish shared governance and ethics

principles for medical AI, enable privacy-preserving federated data collaboration, define standardised validation pathways, strengthen human capacity and trust, embed equity and generalisability as core deployment criteria, and build sustained joint research and innovation ecosystems. Indo-French cooperation in health can thus serve as a practical testbed for globally relevant, responsible medical AI deployment.

AI & Automotive

Artificial intelligence is becoming a central driver of transformation across the automotive value chain, accelerating the shift toward software-defined mobility. It is reshaping vehicle intelligence, manufacturing, customer experience, and mobility services while redefining competitiveness and industrial sovereignty. Within the Horizon 2047 partnership, France and India share an interest in moving beyond adoption toward greater AI sovereignty in mobility. Their complementarities—industrial engineering and regulatory maturity on the French side, scale, digital talent, and data ecosystems on the Indian side—create favourable conditions for a joint positioning as a credible “third pole” in automotive AI.

Despite strong potential, several obstacles persist: difficulty scaling AI beyond proof-of-concept phases, legacy industrial systems and hardware constraints, fragmented standards and data frameworks, and data sovereignty and privacy restrictions.

The paper recommends coordinated policy and industrial initiatives, including a Franco-Indian Automotive AI Mission, a trusted mobility data corridor based on federated learning, shared sovereign compute access, harmonised safety and certification approaches, support for semiconductor and edge-AI capabilities, and strengthened supply-chain resilience. These measures aim to accelerate industrial deployment while reinforcing competitiveness, safety, and sustainability.

AI & Regulation

Given the strategic importance of trust and interoperability, the white paper also highlights the need for structured dialogue on AI governance and regulatory approaches. France and India operate within distinct but increasingly influential regulatory environments, shaped respectively by European frameworks and India’s rapidly evolving digital governance architecture.

Rather than seeking regulatory convergence, the objective is to clarify points of alignment and friction in areas such as data governance, standards, certification, transparency, and accountability. Greater mutual understanding can facilitate cross-border innovation while ensuring compatibility with national and regional legal frameworks.

The paper recommends establishing a sustained Indo-French dialogue on AI governance and standards, promoting interoperability and trusted data collaboration mechanisms, sharing best practices on evaluation and certification of high-risk AI systems, and encouraging joint contributions to international standard-setting and governance discussions.

A first step toward structured cooperation

As a first document, this white paper provides a shared baseline for action. It consolidates initial observations, identifies operational challenges, and proposes pragmatic recommendations that can be tested, refined, and expanded over time.

By connecting scientific cooperation, industrial deployment, and governance dialogue, the France-India AI Initiative seeks to progressively build a durable framework for responsible, inclusive, and sustainable AI development—grounded in trust, measurable impact, and long-term partnership. The France India AI Initiative white paper is endorsed by Anne Bouverot, France’s Special Envoy for AI, Chair of the Board of Ecole Normale Supérieure.

Key Recommendations

AI & Healthcare recommendations

- Establish a Shared Indo-French Governance and Ethics Framework for Medical AI.
- Enable Privacy-Preserving Data Collaboration Through Federated Approaches.
- Define Standardized Translational and Validation Pathways for Medical AI.
- Strengthen Human Capacity, Training, and Trust Within Health Systems.
- Embed Equity and Generalizability as Core Criteria for Deployment.
- Build Sustained Scientific, Educational, and Innovation Ecosystems.

AI & Automotive recommendations

- Launch a dedicated Franco-Indian Automotive AI Mission.
- Create shared Indo-French Mobility Data Corridor and Infrastructure Access.
 - Project "Setu-Pont" : A legally compliant, encrypted Indo-French data pipeline.
 - Sharing AI infrastructure with a Sovereign Compute Access Pass.
 - "Sadak-Rue" : The Open Edge-Case Dataset.
- Promote Mutual Recognition of AI Safety Certification with Tax Incentives for Safety & Cybersecurity.
- Incentivize Edge AI and Semiconductor Development.
- Leverage GCCs as Innovation Hubs.

- Promote Standardization of the Industrial Metaverse.
- Strengthen the Supply Chain with and for AI.
 - The SME AI Uplift (Tier 2/3 Modernization) and Co-creation among OEMs & Startups.
 - 'Mobile Battery' Trading License.
 - Semiconductor supply resilience.
- Establish Indo-French Center of Excellence in Automotive AI.
- Large-Scale Skilling Programs.
 - Student Up-skilling and Cross-skilling Programs.
 - Workforce Reskilling Programs.
 - Executive AI Ethics & Management Exchange.
- Joint Internships and Fellowships linking student talent pipeline with industry.

AI & Regulation recommendations

- Promote Franco-Indian alliance as a driver of normative convergence.
- Promote Practical, Ethical, and Beneficial National AI Use Cases.
- Establish a joint AI regulatory sandbox to enable collaborative testing of innovative AI.
- Implement a coordinated, regulator-led approach for context-specific oversight.
- Maintain a technology-neutral, future-proof approach.
- Foster partnerships between private-sector companies.
- Create a strategic alternative through a strong AI partnership between France and India.
- Promote a global grammar for trustworthy AI.

Context & Introduction



Alexandre Mirlesse

Diplomat, Envoy for Artificial Intelligence and New Technologies, Africa–France Summit (Nairobi), Ministry for Europe and Foreign Affairs, France, Young Leader.



Arthur Barichard

Director,
Artificial Intelligence and
Digital Council, France.

Artificial intelligence is now a major driver of transformation. It is accelerating scientific innovation, reshaping entire industrial sectors, altering economic models, and increasingly influencing societies' ability to deliver public and private services that are more efficient, more inclusive, and more resilient. As these technologies diffuse, the associated challenges become clearer: access to data and computing power; the quality, evaluation, and validation of solutions; skills and adoption; security and information integrity; environmental sustainability; as well as market structure, competition conditions, and risks of concentration.

In this context, the Franco-Indian relationship offers distinctive potential. It rests on structural complementarities—scientific excellence, industrialisation capacity, market depth, and diversity of use cases—together with a growing willingness to organise concrete cooperation. It also sits at the crossroads of several international dynamics: the proliferation of forums dedicated to AI governance, intensifying technological competition, the rapid evolution of regulatory frameworks, and rising societal expectations around transparency, trust, and measurable impact.

A diplomatic moment: from awareness to collective action

Gathered in Paris on 10–11 February 2025, participants from more than 100 countries—states, international organisations, civil society, the private sector, academia, and the research community—recognised that the rapid development of AI technologies is bringing about a major paradigm shift, calling for responses that are both ambitious and pragmatic. The declaration adopted at the conclusion of the AI Action Summit set out a clear framework: to act in the service of the public interest and help close the digital

divide, drawing on three guiding principles—science, solutions (including an emphasis on open AI models that respect national frameworks), and standards, in accordance with international frameworks. The declaration also highlighted key priorities: strengthening ecosystem diversity; promoting more inclusive AI; encouraging deployment that is ethical, safe, secure, trustworthy, and human-centred; preventing excessive market concentration; supporting AI that is sustainable for people and the planet; and reinforcing international cooperation as well as coordination of AI governance. Taken together, these orientations reflect a shared conviction: AI cannot be approached solely as a technical matter. It requires trust frameworks, robust validation methods, interoperability mechanisms, and sustained venues for dialogue capable of linking principles to implementation.

It is within this landscape that the Franco-Indian dynamic has been strengthening. It is not intended to replace existing multilateral frameworks; rather, it can contribute to them—through enhanced partnerships, the development of operational cooperation frameworks, the connection of relevant stakeholders, and, more broadly, a cooperation approach oriented towards tangible and verifiable results.

France–India: complementarities conducive to structured cooperation

Franco-Indian cooperation on AI is developing from a combination of distinctive strengths whose interaction is particularly promising.

France benefits from a leading research ecosystem, recognised expertise in governance, evaluation, and ethics, and high-value industrial sectors (health, mobility, energy, industry, services). It also draws on extensive experience in standardisation and regulation—an important lever for trust, social acceptability, and the diffusion of solutions.

India stands out for its capacity to deploy at scale, a uniquely deep pool of engineering and technical talent, and a momentum in public digital infrastructure that profoundly reshapes the conditions for technology adoption and digital autonomy. This trajectory—illustrated in particular by India Stack—shows that it is possible to create environments conducive to scaling, while also revealing specific challenges: heterogeneous contexts, data quality, integration into existing systems, and heightened requirements around security and fairness.

France and India also share mutual leadership in the global AI governance landscape. From the Global Partnership on AI to the AI Action Summit, both countries have worked hand in hand in order to build a more robust, inclusive and efficient AI governance. Now is a pivotal time in this regards, France holds the 2026 G7 presidency while India holds the BRICS' one.

These complementarities create a natural space for cooperation. They make it possible to connect innovation with implementation; trust frameworks with impact; research with industrialisation. They also invite an explicit focus on key areas of attention: interoperability, validation, cybersecurity, compliance

with legal frameworks, data governance, transparency and bias mitigation, and the sustainability of infrastructures and models.

Shared challenges: trust, inclusion, sustainability, competitiveness

Like elsewhere, Franco-Indian cooperation on AI faces a recurring paradox: the opportunities are significant, but the real value is created—or lost—in execution. Several structural challenges repeatedly emerge in feedback from public and private stakeholders, researchers, and practitioners:

Data and conditions of access

AI—particularly for high-impact applications (health, mobility, industry)—depends on the availability, quality, and governance of data. The issue is not only access, but also standardisation, documentation, auditability, security, and the ability to organise cross-border collaborations compatible with legal frameworks.

Validation, evaluation, and accountability

AI systems applied to sensitive domains require robust evaluation methodologies: performance, robustness, generalisability, bias, explainability, as well as monitoring and update conditions. The challenge is to build validation frameworks that are proportionate to risk, adapted to real-world contexts of use, and compatible with accountability requirements.

Trust, security, and information integrity

Social acceptance of AI depends on transparency, security, protection against malicious uses, and the ability to address risks related to information integrity. This dimension is now inseparable from competitiveness: without trust, deployments slow; without security, negative impacts can become systemic.

Skills, adoption, and organisational transformation

AI is not a simple “add-on”: it reshapes processes, jobs, and the organisation of work. Training, change management, internal governance adaptation, and the development of steering capabilities (technical, legal, ethical, and domain-specific) are critical success factors—and often underestimated.

Environmental sustainability and infrastructure

Infrastructure, compute, energy, and model efficiency are becoming central issues. They require trade-offs and innovation (hardware, architectures, energy efficiency) that must be integrated from the design phase, not addressed only after deployment.

These challenges strongly echo the spirit of the Paris declaration: science (quality, rigour, evaluation), solutions (use cases, impact, scaling), and standards (trust, interoperability, governance). They provide a relevant framework to progressively structure a Franco-Indian cooperation agenda capable of producing tangible outcomes.

The France India AI Initiative: a collegial approach driven by Young Leaders

The France India AI Initiative was designed as a response to this need for method and continuity. Led by the Fondation France-Asie and the France India Foundation, with a decisive mobilisation of their Young Leaders, it draws on a diversity of expertise—researchers, companies, practitioners, and civil society representatives. Its aim is to provide a collective working framework that enables:

- sharing diagnoses and observations grounded in practitioners’ experience;
- identifying common priorities and pragmatic cooperation opportunities;
- formulating actionable recommendations that can be tested and refined;
- aggregating, over time, additional contributors and new themes.

The emphasis on Young Leaders is not incidental: it reflects an operational conviction. AI-driven transformations play out not only in laboratories and major institutions, but also across organisations, value chains, hospitals, industrial companies, and public administrations. Young Leaders—through their professional anchoring, their ability to bridge disciplines and sectors, and their proximity to execution constraints—are well positioned to identify concrete conditions for success as well as friction points that must be addressed from the outset.

The initiative’s approach is therefore intended to complement institutional frameworks: it seeks to create a space for dialogue and proposal-building that connects scientific, industrial, and regulatory levels and maintains continuity over time, beyond event-driven sequences.

A first “working paper”: three workstreams to structure cooperation

This white paper should be read for what it is: a first white paper—an initial, structuring yet imperfectible baseline intended to evolve and deepen over time. It proposes a first cycle organised around three workstreams:

AI & Healthcare

Exploring the conditions for strengthened cooperation in health: data access, clinical validation, bias mitigation, professional adoption, and the link between research and deployment. Health is a natural area for cooperation, both because potential impact is high and because requirements around trust, security, and accountability are particularly demanding.

AI & Automotive

Analysing levers for innovation and competitiveness across a value chain undergoing rapid transformation (software-defined vehicle, ADAS, production, supply chain, quality). Automotive illustrates a central challenge: connecting technological transformation, industrialisation, standardisation, and industrial sovereignty in a highly globalised sector.

AI & Regulation

Providing a comparative overview and key points of attention regarding governance frameworks: regulatory approaches, standardisation, transparency requirements, and conditions for trust. This workstream is more transversal in nature: it seeks to clarify areas of convergence and friction so that cooperation remains compatible with each party's frameworks and responsibilities.

These three themes were selected as practical entry points—two application verticals (health, mobility) and one transversal axis (governance)—covering a broad spectrum from research and use cases to trust and interoperability conditions.

Towards a lasting platform: map, connect, experiment, expand

Beyond this first cycle, the ambition of the France India AI Initiative is to become a lasting platform for exchange and dialogue, capable of progressively bringing in new actors and new topics. This ambition rests on a simple logic: creating collective value through four complementary movements.

Map

Clarify the places, programmes, initiatives, platforms, and stakeholders involved in Franco-Indian AI cooperation: research, industry, startups, public administrations, hospitals, standardisation bodies, regulators, and infrastructure providers. A shared mapping helps reduce fragmentation, identify bridges, and make opportunities more visible.

Connect

Bring the right expertise together, at the right level, around concrete objectives: applied research projects, benchmarks, validation protocols, systems interoperability, standards, and training programmes. Cooperation strengthens when stakeholders have a stable dialogue framework and clearly identified points of contact.

Experiment

Promote testable cooperation: pilots, demonstrators, “sandboxes,” evaluation protocols, and industrial partnerships. The goal is to narrow the gap between intention and execution and generate shareable lessons learned.

Expand

Open the platform to additional themes over time, based on identified priorities: AI and energy, AI and skills, trust infrastructure, related critical technologies, and more. The objective is to avoid a one-off effect and build a durable framework that can evolve in step with emerging challenges.

In this spirit, the present white paper is a foundational step: it consolidates an initial diagnosis, identifies operational obstacles, and proposes recommendations to structure Franco-Indian cooperation. It also seeks to establish a method—linking science, solutions, and standards within an open, multi-stakeholder dialogue oriented towards the public interest.

An open, progressive, trust-oriented approach

The France India AI Initiative is grounded in a constructive, progressive, and open approach. It recognises the diversity of contexts, frameworks, and priorities, and favours a cooperation trajectory built on trust, methodological rigour, and operational usefulness. The objective is not to set out a fixed vision, but to provide a working baseline that can be discussed, refined, and strengthened—as the platform welcomes new contributors and as cooperation efforts translate into concrete outcomes.

It is on these terms that Franco-Indian cooperation on AI can continue to deepen: not only through dialogue, but also through projects, standards, validation frameworks, and trust mechanisms capable of supporting responsible, inclusive, safe, secure, and sustainable deployment—for the benefit of societies, economies, and the public interest.

THE FRANCE INDIA AI INITIATIVE

Artificial intelligence (AI) is poised to redefine the contours of the global economy, scientific research and technological innovations. While the United States and China compete for dominance in this sector, India, through its talent resource, and France, through its R&D, have major assets to compete in this race. In this context, the France India Foundation, the Indian chapter of the Fondation France-Asie, is launching its France India AI Initiative, an ambitious program aimed at strengthening cooperation between France and India in the field of AI.

The France India AI Initiative reflects a shared ambition to position artificial intelligence as a cornerstone of the strategic partnership between France and India. Led by the France India Foundation and the Fondation France-Asie, this initiative brings together the collective intelligence of more than 60 experts from both countries and leading experts from academia, industry, healthcare, and public institutions. The objective is clear: to identify actionable ideas and policy-oriented recommendations that can transform complementary national strengths into a durable and impactful Franco-Indian AI collaboration.

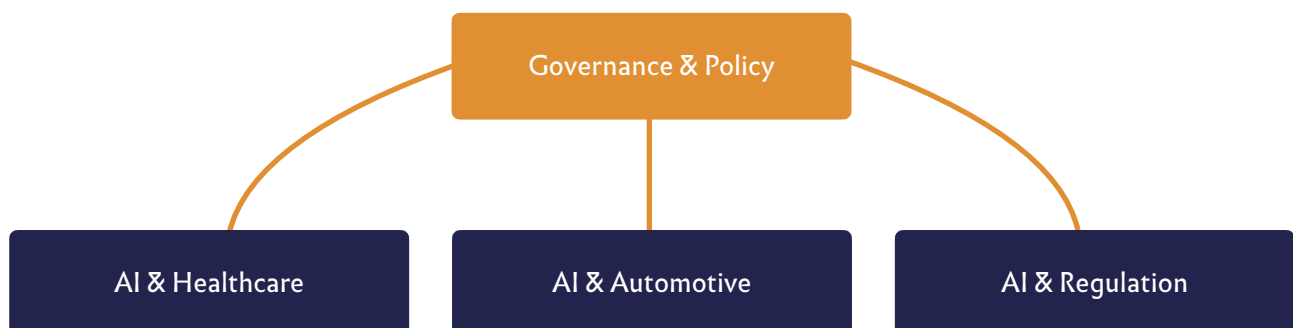
France and India occupy highly synergistic positions in the global AI landscape. France

brings long-standing excellence in fundamental research, ethics-by-design, regulation, and high-value industrial systems. India contributes unmatched scale, engineering talent, data diversity, and the capacity to deploy AI solutions at population and market scale.

This white paper gathers the first recommendations and articulates how these strengths can be strategically aligned to address shared challenges and global priorities. The document focuses on three priority domains where Franco-Indian collaboration can generate immediate and long-term impact:

- AI & Healthcare,
- AI & Automotive,
- and AI & Regulation.

3 working groups to shape policy-oriented first recommendations.



The France India AI Initiative aims at serving as a model for global cooperation, moving AI discussions beyond a Western-centric framework toward a more inclusive and diverse dialogue. By fostering collaboration between two nations with strong technological capabilities and distinct regulatory perspectives, this initiative and first version of the white paper could help shape a balanced and ethical approach to AI governance.

Moreover, it could play a key role in promoting AI adoption and literacy, particularly in France, where public skepticism toward AI remains a challenge. The last AI Action Summit that took place in Paris in

February 2025 reminded us that AI should not be framed purely as a risk-mitigation topic but as a means for innovation, economic growth, and public service improvement.

Inspired by successful cooperation programmes such as the France China Climate Initiative of its sister foundation, the France China Foundation—the France India AI Initiative aims to facilitate exchanges between entrepreneurs, researchers, academic institutions and public actors. It will help foster dialogue and collaboration to identify common opportunities and address key challenges related to artificial intelligence.

4 pillars of the France India AI Initiative

Developing Research
& Innovation

Strengthening Scientific
& International Cooperation

Facilitating Engagement between AI
Stakeholders

Promoting AI for all
& for Public Good

Milestones & Contributors

MILESTONES OF THE INITIATIVE

October, 9th 2025 – Roundtable in India.

During the 2025 Young Leaders seminar in India, a roundtable was held to highlight different perspectives on innovation, ethics, inclusion, and opportunities related to the emergence of AI on a global scale.

Keynote experts (by alphabetical order):

Flore Cousin, Co-founder & COO, The Marshmallow Project

Sarita Kaloya, Senior Director, Capgemini

Pratyush Kumar, Co-Founder & CEO, Sarvam AI and AI4Bharat

Alexandre Mirlesse, Diplomat, Envoy for Artificial Intelligence and New Technologies, Africa–France Summit (Nairobi), Ministry for Europe and Foreign Affairs, France

Charlie Perreau, Head of Tech–Media Start-up Department, Les Echos

Aakrit Vaish, Co-Founder, Activate & Founder, Haptik

October, 14th 2025 – Roundtable in France.

With the support of Brunswick Group, diplomats, researchers, and experts gathered at Albert School to open a strategic dialogue on artificial intelligence between France and India in light of American and global developments.

Keynote experts (by alphabetical order):

Ahmed Baladi, Partner, Gibson Dunn, Young Leader

Arthur Barichard, Director, Artificial Intelligence and Digital Council, France, Former Deputy Ambassador for Digital Affairs and Artificial Intelligence, France.

Philippe Cordier, Chief AI Scientist and VP Data Science, AI and Data Engineer, Capgemini Invent (partenaire de la Fondation France–Asie)

Alexandre Escargueil, PhD in molecular and cellular biology, Professor, Sorbonne University

Michael Fitzpatrick, Partner, Brunswick Group

Rahul Gaurav, Neuroscientist, AI–MRI, Researcher, Paris Brain Institute

Mark Seifert, Partner, Brunswick, Co–Head of Cyber Practice

December 2, 2025 – Official Pre-Summit Event of the AI Impact Summit 2026.

Addressing the complex challenges of international, cross-continental research and academic partnerships, the initiative organised an Indo-French Dialogue on AI in Healthcare: Ethics, Data, Challenges, and Opportunities. This was supported by the Fondation France-Asie, Sorbonne University and Paris Brain Institute and labeled by India AI as an Official Pre-Summit Event of the AI Impact Summit 2026.

Keynote experts (by alphabetical order):

Charles Assaad, Researcher, Inserm

Ségolène Aymé, Geneticist, Founder, Orphanet

Manon Bachy, Professor, Pediatric Orthopedic Surgeon, AP-HP Trousseau

Manidipa Banerjee, Professor and Scientific Coordinator of Indo-French Campus, IIT Delhi

Eléonore Bayen, Professor of Physical Medicine and Rehabilitation, Sorbonne Université

Jean-Cassien Billier, Lecturer in Moral and Political Philosophy, Sorbonne University

Sylvie Bothorel, Chief Data Officer, Paris Brain Institute

Imen Canova, PhD, Data Science Senior Manager, BioMérieux

Olivier Colliot, Research Director, Center for AI and Data Science, Paris Brain Institute

Rajinder K Dhamija, Chair, National Task Force on Brain Health of Government of India; Director, Institute of Human Behaviour and Allied Sciences

Laurent Drazek, PhD, Innovation Leader R&D Data Science Expert, BioMérieux

Guillaume Fiquet, Vice-president, International and Economic Partnerships, Sorbonne Université

Rahul Gaurav, Neuroscientist, AI-MRI, Researcher, Paris Brain Institute

Edith Gross, International Scientific Affairs Manager, Paris Brain Institute

Corinne Isnard, Professor & Researcher, Digital Health Expert, Sorbonne Université

Sonia-Athina Karabina, Professor, Sorbonne Université, Researcher, Sorbonne Center for Artificial Intelligence

Srini Kaveri, Director Maison de l'Inde, Former Director CNRS India

Guillaume Klossa, CEO, T Life

Sandeep Kumar, Associate Professor, Department of Electrical Engineering, School of Artificial Intelligence, IIT Delhi

Nand Kumar, Prof. In charge ICMR CARE in Neuromodulation for Mental Health, AIIMS Delhi

Senthil Kumaran, Senior Professor and Faculty in charge, Neuroimaging, AIIMS Delhi

Brian Lau, Scientific Director, Paris Brain Institute

Charles Mahy, Counsellor for Health and Social Policies, Embassy of France to India

Maria Melchior, Research Director, Inserm

Alexandre Mirlesse, Diplomat, Envoy for Artificial Intelligence and New Technologies, Africa-France Summit (Nairobi), Ministry for Europe and Foreign Affairs, France

Augustin Missoffe, CEO Mage-X

Pierre Pouget, President Ethical Committee, Professor & Researcher, Paris Brain Institute

Shalia Shah, Counsellor, Embassy of India to France

M Srinivas, Director, All India Institute of Medical Sciences (AIIMS), New Delhi

Antoine Tesnière, CEO, ParisSanté Campus

Marc Thellier, MSc-Ph in the Parasitology Department, APHP

Franck Verdonk, Professor of Medicine, Co-founder SurgeCare, Sorbonne Université

Raphaël Vialle, Vice Dean for International Affairs and Europe, Sorbonne University; Pediatric Orthopedic Surgeon, AP-HP Trousseau

Marie Vidailhet, Professor of Neurology, Pitié Salpêtrière Hospital Sorbonne APHP

CONTRIBUTORS OF THE INITIATIVE

Young Leaders (by alphabetical order)

Vishal Agrawal, Portfolio Manager – Equities, Axis Asset Management Company

Laurie-Anne Ancenys, Partner, Head of Tech & Data Practice, Paris Office, A&O Shearman

Neha Arolkar, Director, Capgemini

Manon Bachy, Professor, Pediatric Orthopedic Surgeon, AP-HP Trousseau

Ahmed Baladi, Partner, Gibson Dunn

Flore Cousin, Co-founder & COO, The Marshmallow Project

Thibaud Frossard, Chief of Staff to the CEO of the Capgemini Group

Grégoire Genest, Founder & CEO, Albert School

Tatiana Jama, Founder & Managing Director, SISTAFUND, Co-Founder, SISTA

Sarita Kaloya, Senior Director, Capgemini

Vanshica Kant, Associate Investment Officer, Public Sector Clients, Financial Institutions and Funds Department at the Asian Infrastructure Investment Bank

Arjun Kothari, Managing Director, Kothari Sugars and Chemicals Ltd, Kothari Petrochemicals Ltd

Pratyush Kumar, Co-Founder & CEO, Sarvam AI and AI4Bharat

Alexandre Mirlesse, Diplomat, Envoy for Artificial Intelligence and New Technologies, Africa–France Summit (Nairobi), Ministry for Europe and Foreign Affairs, France

Augustin Missoffe, CEO, Mage-X

Charlie Perreau, Head of Tech-Media Startup Department, Les Echos

Abhinav Prakash, Singh National Vice-President, BJP Youth Wing, University of Delhi

Harshad Reddy, Director, Group Oncology & International, Apollo Hospitals

Robin Rivaton, Entrepreneur and author, CEO, Stonal

Antoine Tesnière, Chief Executive Officer, ParisSanté Campus

Chinmay Tumbe, Faculty Member, Economics Area, Indian Institute of Management, Ahmedabad

Aakrit Vaish, Co-Founder, Activate & Founder, Haptik

Raphaël Vialle, Chairman of Pediatric Orthopedic and Surgery, Hôpital Trousseau AP-HP

Other contributors (by alphabetical order)

Charles Assaad, Researcher, Inserm

Sékolène Aymé, Geneticist, Founder, Orphanet

Manidipa Banerjee, Professor and Scientific Coordinator of Indo–French Campus, IIT Delhi

Arthur Barichard, Director, Artificial Intelligence and Digital Council, France

Eléonore Bayen, Professor of Physical Medicine and Rehabilitation, Sorbonne Université

Priyanka Bhat, All India Institute of Medical Sciences (AIIMS) Delhi

Jean-Cassien Billier, Lecturer in Moral and Political Philosophy, Sorbonne University

Sylvie Bothorel, Chief Data Officer, Paris Brain Institute

Imen Canova, PhD, Data Science Senior Manager, BioMérieux

Olivier Colliot, Research Director, Center for AI and Data Science Paris Brain Institute

Philippe Cordier, Chief AI Scientist and VP Data Science, AI and Data Engineer, Capgemini Invent

Stéphanie Debette, Director, Paris Brain Institute

Rajinder K Dhamija, Chair, National Task Force on Brain Health of Government of India Director, Institute of Human Behaviour and Allied Sciences



Laurent Drazek, PhD, Innovation Leader R&D Data Science Expert, BioMérieux

Sunu Engineer, Chief Technology Officer at Ryussi Technologies (P) Ltd.

Alexandre Escargueil, PhD in molecular and cellular biology, professor at Sorbonne University

Guillaume Fiquet, Vice-president, International and Economic Partnerships, Sorbonne Université

Michael Fitzpatrick, Partner, Brunswick Group

Xavier Fresquet, Deputy Director, Sorbonne Center for Artificial Intelligence, Sorbonne University

Rahul Gaurav, Neuroscientist, AI-MRI , Researcher Paris Brain Institute

Michelle George, Paris Brain Institute

Edith Gross, International Scientific Affairs Manager, Paris Brain Institute

Corinne Isnard, Professor & Researcher, Digital Health Expert, Sorbonne Université

Magali Jaillard Dancette, Data Scientist, Sequencing Department, bioMérieux

Saumya Jetley, Faculty, Plaksha University Ph.D. in Computer Vision and Machine Learning Algorithms, University of Oxford

Harshit Kacholiya, Tech Policy and Strategy at iSPIRT

Sonia-Athina Karabina, Professor, Sorbonne Université, Researcher, Sorbonne Center for Artificial Intelligence

Srini Kaveri, Director Maison de l'Inde, Former Director CNRS India

Guillaume Klossa, CEO, T Life

Nand Kumar, Prof. In charge ICMR CARE in Neuromodulation for Mental Health, AIIMS Delhi

Sandeep Kumar, Associate Professor, Department of Electrical Engineering, School of Artificial Intelligence, IIT Delhi

Senthil Kumaran, Senior Professor and Faculty in charge , Neuroimaging, AIIMS Delhi

Brian Lau, Scientific Director, Paris Brain Institute

Frédérique Lesaulnier, Data Protection Officer, Paris Brain Institute

Charles Mahy, Counsellor for Health and Social Policies, Embassy of France to India

Maria Melchior, Research Director, Inserm

Pierre Pouget, President Ethical Committee, Professor & Researcher, Paris Brain Institute

Juan-Fernando Ramirez, Director, Global Health Institute of Alliance Sorbonne University

Mark Seifert, Partner, Brunswick, Co-Head of Cyber Practice

Shalia Shah, Counsellor, Embassy of India to France

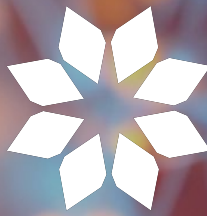
M Srinivas, Director All India Institute of Medical Sciences (AIIMS), New Delhi

Shreya Srivastava, Paris Brain Institute

Marc Thellier, MSc-Ph in the Parasitology Department, APHP

Franck Verdonk, Professor of Medecine, Co-founder SurgeCare, Sorbonne Université

Marie Vidailhet, Professor of Neurology, Pitié Salpetriere Hospital Sorbonne APHP



FRANCE INDIA
FOUNDATION

AI & Healthcare

France India AI Initiative
February 2026

AI & Healthcare

Heads of AI & Healthcare working group



Raphaël Vialle

Chairman of Pediatric
Orthopedic and Surgery,
Hôpital Trousseau AP-HP,
Young Leader.



Antoine Tesnière

Chief Executive Officer,
PariSanté Campus,
Young Leader.



Rahul Gaurav

Neuroscientist, AI-MRI
Researcher, Paris Brain
Institute.

This white paper section sets out three core policy conclusions:

- Artificial intelligence (AI) in health should be treated as public-interest infrastructure rather than a technology race, if it is to genuinely support human well-being, trust, and long-term sustainability.
- Brain health represents the most demanding and governance-sensitive domain for medical AI. As such, it provides a credible and appropriate flagship for Indo-French cooperation before extending these approaches to wider global health applications.
- India and France are particularly well placed to translate trustworthy medical AI into practice within the next 12–36 months, building on existing legal and institutional frameworks through coordinated action on governance, validation, data collaboration, and human capacity.

The India AI Impact Summit 2026 has articulated a vision for international cooperation structured around seven interconnected pillars, referred to as chakras, emphasizing human capacity, social inclusion, trust, resilience, scientific progress, and economic and societal benefit. Taken together, these pillars reflect a broader policy ambition: ensuring that emerging technologies contribute meaningfully to human well-being, long-term sustainability, and shared prosperity. Moreover, the India AI Summit is structured around three foundational principles expressed as sutras, which frame policy action in terms of People, Planet, and Progress. They provide a reference for translating strategic objectives into practical policy measures and for aligning technological development with societal priorities, environmental considerations, and sustainable economic growth.

This Indo-French white paper is grounded in that ambition and draws inspiration from the principles articulated under the India AI Impact Summit framework and sets out a health-first, internationally oriented scientific approach, with a specific focus on brain health and global healthcare systems. Brain health is used as a flagship policy domain because it represents the most stringent test case for responsible medical AI. While the paper examines how advances in data-driven medicine and computational methods can be applied to strengthen health systems, support earlier and more reliable diagnosis, and improve equity within the Indo-French context, its analysis is also applicable to broader challenges in access to care across diverse populations, with particular attention to brain health and complex health conditions. The recommendations derived from this analysis are intended to guide policy development and bilateral cooperation by identifying priority areas for action, practical implementation pathways, and mechanisms for sustained scientific collaboration in global healthcare. This paper does not treat AI as an end in itself, but as an enabling capability whose public value depends on governance, scientific validation, and conditions of real-world deployment. It therefore focuses on practical, implementation-oriented pathways for global health cooperation and the responsible use of emerging technologies.

Furthermore, it positions Indo-French cooperation not as a symbolic partnership, but as a practical laboratory for the global governance of medical AI. France brings a rich and mature health data ecosystem, built on extensive health data assets, including hospital data warehouses (EDS), large-scale cohorts, and the National Health Data System (SNDS), and supported by leading academic and hospital stakeholders (Inserm, AP-HP). This ecosystem operates within a robust and protective legal framework, including the GDPR and the AI Act, and is complemented by administrative simplification mechanisms, notably reference methodologies, that facilitate responsible access to data. It is underpinned by well-established authorities and governance structures, including the CNIL and the Health Data Hub. India, in turn, brings integrated digital health infrastructures (Ayushman Bharat, the Ayushman Bharat Digital Mission, and the AIIMS network), as well as scale and population diversity. Together, they form a complementary ecosystem capable of supporting the testing, validation, and responsible scaling of trustworthy medical AI in France, India, and other parts of the world.

1. Strategic Review and Priorities

Lack of a Predictable Regulatory and Validation Pathway for Medical AI

Medical AI systems currently lack a governance, validation, and translational pathway comparable to drug or medical device approval processes. There is no shared framework to evaluate safety, efficacy, robustness, and security before deployment across clinical settings. This creates uncertainty for regulators, clinicians, and patients, and limits responsible scaling of AI in healthcare.

Data Access and Privacy Constraints Limiting Cross-Border Collaboration

Strict data protection and privacy regulations have created structural barriers to cross-border collaboration in health research. In the absence of trusted mechanisms for lawful collaboration, AI models remain trained on limited and non-representative datasets, constraining generalizability and clinical relevance, particularly in international contexts such as Indo-French cooperation.

Persistent Translation Gap from Research to Clinical Practice

Despite significant advances in medical AI research, many models fail to progress beyond experimental stages. The lack of standardized translational pipelines from methodological research to clinical validation and deployment results in limited real-world impact and underutilization of public research investments.

Trust, Skills, and Adoption Deficits within Health Systems

Health professionals and end users may lack sufficient opportunities for training in the usability, explainability, and limitations of AI systems. This can contribute to low trust, inconsistent adoption, and resistance to deployment, even where technical performance is strong. Human capacity and trust have emerged as critical bottlenecks for implementation.

Risk of Bias and Inequity at Population Scale

AI systems trained on non-diverse or incomplete datasets risk reinforcing existing health disparities across socio-economic, linguistic, and geographic groups. Without deliberate attention to data diversity and validation across populations, large-scale deployment of medical AI may exacerbate inequities rather than reduce them.

Fragmentation of Long-Term Scientific and Educational Collaboration Frameworks

While individual collaborations and mobility initiatives exist, they remain fragmented and short-term. The absence of sustained, structured frameworks linking education, research, validation, and innovation limits continuity, knowledge transfer, and the long-term impact of Indo-French cooperation in AI for health.

2. Scope

The scope of our work is intentionally limited to policy and implementation considerations related to the use of AI in regulated healthcare settings, with brain health serving as a flagship domain for Indo-French cooperation, due to its reliance on complex multimodal data, long-term care trajectories, and high societal and economic impact.

This paper does not assess or compare specific AI technologies, algorithms, datasets, or commercial products, nor does it provide technical, clinical, or performance evaluations. It does not propose new regulatory instruments, amend existing legal frameworks, or replace applicable national, European, or international regulations. Furthermore, non-medical applications of AI, including consumer-facing systems, defense or security uses, and applications outside regulated healthcare environments, fall outside the scope of this analysis.

Taken together, this paper aims to inform policy dialogue and bilateral cooperation between India and France. It does not constitute regulatory guidance, technical standards, or operational protocols.

3. Observations

The Indo-French Dialogue on AI in Healthcare, held on December 2nd, 2025, was organized with the support of Sorbonne University, the Paris Brain Institute, the Indian Institute of Technology (IIT), the All India Institute of Medical Sciences (AIIMS), and the Fondation France-Asie. The discussions brought together perspectives from research, clinical practice, public institutions, and policy-oriented stakeholders.

Across sessions, there was a shared recognition that advances in medical AI are progressing more rapidly than the capacity of existing health systems and regulatory frameworks to govern, validate, and deploy them at scale. While technical innovation has been substantial, it became clear that institutional, regulatory, and human factors represent the primary constraints on real-world impact.

A recurring observation concerned the absence of a coherent and harmonized governance and validation pathway for medical AI systems. Unlike pharmaceuticals or medical devices, AI-based tools often move through fragmented and inconsistent evaluation processes creating uncertainty regarding safety, efficacy, robustness, and accountability prior to clinical deployment. This uncertainty contributes both to cautious underutilization and, in some cases, premature or uncoordinated implementation.

It was also observed that data protection and privacy regulations, although essential, currently limit the ability to conduct cross-border research using diverse and representative health data. As a result, many AI models are developed and validated on restricted datasets that do not adequately reflect population heterogeneity, constraining generalizability and clinical relevance, particularly in international collaborations such as those between India and France.

Another consistent observation related to the persistent gap between research and routine clinical practice. Despite promising results in controlled research environments, many AI models fail to progress into clinical use due to the absence of standardized translational pathways, multicenter validation mechanisms, and post-deployment evaluation frameworks. This limits the return on public and institutional research investments.

Human capacity and trust emerged as critical determinants of adoption. Health professionals frequently lack sufficient training in the interpretation, usability, and limitations of AI systems, affecting confidence and utilization even where technical performance is strong. It became evident that explainability, workflow integration, and user-centered design are often undervalued relative to algorithmic performance, limiting effective uptake in clinical settings.

Finally, concerns were raised regarding equity and long-term sustainability. Without deliberate attention to data diversity, population-level validation, and sustained collaboration frameworks, the deployment of medical AI risks reinforcing existing disparities in access to care. While numerous bilateral and institutional initiatives exist, they were observed to be largely project-based and time-limited, constraining continuity, knowledge transfer, and cumulative impact.

4. Practical illustrations & Use cases

The following use cases illustrate how the issues identified in this first paper manifest in real-world contexts and why coordinated policy action is required. They are not exhaustive and are presented for illustrative purposes only.

4.1. Multicenter Brain Health Research and Clinical Validation

Brain health and neurological conditions represent a domain where AI has demonstrated significant research potential but limited clinical translation. Multicenter studies across hospitals and research institutions require access to diverse datasets, consistent validation protocols, and alignment across regulatory environments. These studies increasingly involve heterogeneous imaging infrastructures, further complicating harmonized validation and cross-site comparability. This makes brain health an ideal testbed for Indo-French regulatory, data, and validation alignment before scaling to other disease areas. In the Indo-French context, differences in data governance frameworks, validation practices, and clinical workflows complicate cross-border collaboration and limit the generalizability of findings. This use case highlights the need for coordinated approaches to governance, validation, and trust in order to move from research outputs to clinically deployable tools.

4.2. Translation of Medical AI from Research to Routine Clinical Practice

Across multiple disease areas, AI models show promising performance in controlled research settings but fail to reach routine clinical use. This use case reflects how gaps in validation, oversight, and institutional readiness continue to impede the integration of AI into routine care delivery. Importantly, validation guidelines need to be adapted to each stage and type of research, from methodological development to translational proof-of-concept studies, to clinical validation. This is often due to the absence of standardized translational pathways, unclear accountability structures, and limited integration into clinical workflows. In both India and France, health systems face challenges in evaluating when and how AI tools are sufficiently mature for deployment.

4.3. Population-Scale Equity and Generalizability in AI-Enabled Healthcare

AI systems trained on narrow or homogeneous datasets risk performing unevenly across populations differentiated by language, geography, socio-economic status, or health system access. This use case underscores how insufficient attention to data diversity and population-level validation can reinforce existing disparities rather than reduce them. In countries with diverse populations such as India, and in international collaborations involving Europe and the Global South, these limitations become particularly visible.

4.4. Capacity Building and Trust Among Health Professionals

The adoption of AI in healthcare is strongly influenced by the confidence and preparedness of clinicians and health system users. This use case demonstrates that gaps in human capacity and trust, rather than technical performance alone, remain a primary constraint on the integration of AI into clinical and public health settings. Limited training in AI interpretability, usability, and limitations contributes to hesitation, inconsistent utilization, and resistance to deployment.

4.5. Cross-Border Collaboration in Public Health and Disease Surveillance

Public health challenges, including both communicable and non-communicable diseases, such as infectious disease monitoring and antimicrobial resistance, increasingly rely on large volumes of heterogeneous data and predictive analytics. This use case illustrates how governance and coordination challenges extend beyond individual clinical applications to population-level public health use. Cross-border collaboration can enhance early detection and response, but is constrained by regulatory, institutional, and infrastructural barriers. In the Indo-French context, aligning data practices, validation standards, and institutional roles remains a challenge. Beyond brain health, similar governance and validation challenges arise in other data-intensive domains of regulated healthcare, such as AI-assisted microscopic diagnostics for infectious, inflammatory, and hematological diseases, where large-scale data acquisition, interpretability, and clinical integration are critical for public health impact.

Use case 1

ArboTracker: A Predictive Framework to Support Arbovirus Outbreak Preparedness.

by BioMérieux

Co-authored by:

Laurent Drazek, PhD, Innovation Leader R&D Data Science Expert, BioMérieux

Imen Canova, PhD, Data Science Senior Manager, BioMérieux

Arboviral diseases such as dengue, chikungunya, and Zika continue to exert significant pressure on public health systems, particularly in regions marked by climatic diversity, demographic expansion, and rapid urbanization. India stands as a clear illustration of this challenge: dengue transmission varies substantially from year to year, with some seasons marked by strong epidemic waves and others showing unusually quiet patterns before a subsequent resurgence. These irregular, shifting dynamics emerge from the interplay of environmental variability, human mobility, vector ecology, and social behavior. As these factors evolve unpredictably, national authorities face increasing difficulty in anticipating future outbreaks and preparing adequate response strategies.

ArboTracker was developed in response to this growing need for anticipation. Rather than relying solely on retrospective surveillance or reactive monitoring, the framework introduces a forward-looking approach designed to project how arbovirus circulation may evolve over the coming months. Its goal is not simply to produce numerical forecasts, but to clarify the range of possible epidemic trajectories and the level of uncertainty associated with them. By transforming diverse and heterogeneous early signals into structured insights, ArboTracker helps institutions shift from a reactive posture to a more proactive mindset, especially in settings where unpredictability has become a defining characteristic of arboviral epidemiology.

The Need for Predictive Intelligence

Traditional surveillance systems are indispensable for identifying outbreaks, yet they capture events only after transmission has already begun. Arboviruses, however, are often driven by climatic and ecological processes that shift weeks or even months before clinical case numbers begin to rise. These upstream signals—changes in temperature, rainfall, humidity, mosquito breeding conditions, and human interactions with their environment—hold substantial predictive value. Anticipating them provides health authorities with critical time to strengthen preparedness, allocate resources, and plan interventions. ArboTracker fills precisely this strategic gap by integrating environmental and contextual information early enough for decision-makers to act before incidence levels rise.

A High-Level Predictive Framework

Despite drawing on advanced artificial intelligence techniques, ArboTracker's conceptual structure remains intentionally simple. The tool integrates three major categories of information that together shape arbovirus dynamics:

1. Environmental tendencies, which reflect evolving climatic conditions that influence mosquito populations. Temperature anomalies, precipitation patterns, and humidity variations directly affect

mosquito breeding, survival, and mobility. These factors are among the earliest signals of changing transmission potential.

2. Contextual indicators, which encompass socio-economic, demographic, geographic, and ecosystem-level characteristics. These include at least population density, ecological suitability, and mobility patterns, all of which modulate exposure risk and determine how outbreaks may spread.
3. Historical patterns, which serve as a benchmark for differentiating expected seasonal rhythms from atypical behaviors. By comparing current signals with past multi-year trends, the model can identify when a season is likely to diverge from traditional patterns.

ArboTracker's forecasting engine synthesizes these inputs to generate plausible scenario ranges. Rather than offering a single deterministic prediction, the system highlights uncertainty intervals, expected trends, and comparisons with previous seasons. Forecast outputs are made accessible through a clear dashboard designed for multiple user groups—from technical analysts to policy leaders. This accessibility is central to the tool's impact: insight is only useful when it can guide timely and informed action.

Insights from the Indian Sub-Continent Use Case

Indian sub-continent's scale and heterogeneity provide a rigorous test for any forecasting model, yet ArboTracker has demonstrated strong performance in this complex environment. During real-time forecasting exercises across the subcontinent, the system successfully anticipated shifts in dengue activity months before traditional surveillance systems captured them. It flagged early signs of rising transmission, detected abnormal seasonal transitions, and identified periods of accelerated viral spread.

These early warnings were later validated by observed epidemiological trends, confirming the model's ability to provide meaningful foresight even under highly dynamic conditions. In several instances, ArboTracker highlighted the likelihood of upcoming surges while reported case numbers were still low, offering health authorities a clearer view of what the coming months might bring. Such predictive intelligence is especially valuable in India, where climatic gradients, population density, and mobility patterns vary widely between states and districts.

Operational Considerations and Future Opportunities

Although ArboTracker's primary mission is to support public health preparedness, its predictive capabilities also provide value to operational actors such as supply-chain planners, diagnostic manufacturers, and emergency response units. Forecasting likely increases or decreases in arbovirus circulation can support more efficient planning for diagnostic reagent supply, distribution strategies, and inventory management. Early insights into expected demand help reduce shortages during epidemic peaks while avoiding unnecessary stock accumulation in low-incidence periods.

Looking ahead, deeper integration of ArboTracker into national and regional surveillance systems presents an important opportunity. Tailored dashboard views for policymakers, automated links with early warning systems, and incorporation into seasonal preparedness cycles could amplify its impact. As climate change and urbanization continue shaping the behavior of arboviruses worldwide, tools like ArboTracker will become increasingly central to modern epidemic intelligence.

Use case 2

From positive blood culture to actionable genomic AST

by BioMérieux

Co-authored by:

Magali Jaillard Dancette, Data Scientist, Sequencing Department, BioMérieux

Bloodstream infections are among the most time-critical medical emergencies, requiring rapid pathogen identification and accurate antimicrobial susceptibility testing (AST) to guide effective therapy and reduce mortality. Conventional phenotypic AST remains the gold standard but takes 18 to 24 hours after blood culture positivity. PCR-based methods offer faster characterization but are restricted to predefined species panels and provide limited resistance information, insufficient for reliably predicting susceptibility. Whole-genome sequencing (WGS) overcomes these constraints by bypassing lengthy culture steps and delivering broad, high-resolution pathogen characterization. When coupled with machine-learning (ML) models, WGS enables accurate genome-based prediction of antibiotic susceptibility. This use-case presents an overview of a diagnostic use case and explains why genomewide, datadriven approaches can bring high medical value by offering fast and comprehensive genomic AST.

Clinical use case: positive blood culture to treatment decision

When a blood culture flags positive, typically within 12-24 hours of incubation depending on organism load, clinicians must balance empiric broadspectrum treatment with the stewardship imperative to deescalate rapidly. Conventional workflows involve Gram stain immediately after positivity, subculture to generate isolated colonies, phenotypic identification and phenotypic AST, often requiring additional overnight hours. This timeline delays optimal therapy.

Several evaluations of directtoWGS methods have shown that sequencing directly from positive blood culture broth, when combined with sample preparation steps that enrich microbial DNA and reduce host DNA, enables pathogen identification in hours, rather than days, and enables genomic AST prediction at the same time. Such approaches have been explored using both shortread and nanopore sequencing technologies. The feasibility study of a nanoporebased workflow on 200 positive blood cultures found that sequencing time required for correct pathogen identification was often less than 10 minutes, especially for monomicrobial samples, demonstrating WGS's potential to dramatically accelerate early clinical decisionmaking (1).

Biological complexity of antibiotic resistance

Predicting antibiotic susceptibility from genomic data is far from straightforward. Bacteria rely on a wide variety of resistance mechanisms which are often multifactorial and complex and obscure the relationship between genomic sequence and phenotypic expression. In *Pseudomonas aeruginosa*, resistance frequently arises from regulatory mutations affecting efflux pumps, or porin loss, mechanisms that are not reliably captured by single resistance markers listed in current databases (2). In

Enterobacterales, β -lactam resistance may involve diverse combinations of ESBL genes, porin mutations, efflux modulation, or changes in gene copy number (3).

These complexities have been highlighted in several studies (1) (3) (2). For instance, an ML evaluation showed that certain annotated porin mutations in *Klebsiella pneumoniae* occurred at high frequency in phenotypically susceptible isolates, markedly reducing the specificity of marker-based approaches (3). Similarly, ciprofloxacin resistance in *E. coli* was associated with multiple gyrase and topoisomerase mutations, yet many of these mutations were also present in susceptible isolates, complicating approaches that rely solely on genomic markers.

Strategies for genomic AST

Once sequencing data is generated, genomic AST seeks to predict whether an organism is susceptible (S), intermediate (I), or resistant (R) to a panel of antibiotics. Two main strategies exist. The first is resistance marker-based prediction, which infers resistance from the detection of known genes or point mutations catalogued in antimicrobial resistance (AMR) databases. The second is genome-wide ML prediction, which trains models on large collections of genomes with paired phenotypic AST results and learns predictive patterns across the entire genome using k-mers or other sequence features (4). Comparisons across clinical isolate datasets show that these approaches differ in performance and reliability (3).

Marker-based prediction

Resistance-marker detection suffers from both false positives and false negatives. False positives occur because the presence of a resistance gene does not always guarantee phenotypic resistance: weak markers may act only in specific genetic backgrounds, require sufficient expression, or depend on combinations of mutations. Large evaluations have shown that resistance markers can appear in a substantial fraction of phenotypically susceptible isolates, leading to overcalling resistance and unnecessarily restricting therapeutic options (3). False negatives arise because AMR databases remain incomplete, thus the absence of known markers cannot ensure susceptibility. Some phenotypically resistant isolates do not carry any recognized markers, potentially causing unsafe false susceptible predictions. Finally, marker-based methods inherently produce binary outputs, lacking the ability to classify intermediate (I) strains, an important clinical category.

ML-based prediction

In contrast, ML models, trained directly on phenotypes, can learn to assign S/I/R categories and better capture predictive signals that extend beyond known AMR determinants. Unlike marker-based methods, ML can identify unrecognized genomic features, model nonlinear interactions across genes and regulatory regions, and maintain accuracy even when resistance databases are incomplete. These ML models are inherently agnostic to known resistance determinants, and can, when trained on datasets that are sufficiently large and representative of genetic diversity, more effectively accommodate the complexity of genome-phenotype relationships.

Clinical evaluations highlight this advantage. In a 956-isolate study, ML achieved 96.9% categorical agreement with phenotypic AST, with very major error rates of only 1 to 2% for qualified models and far greater stability across datasets. In direct comparisons, ML reached 96% accuracy on a binary S/NS task, versus 83.5% for marker-based prediction (3). Feasibility studies using direct-from-blood nanopore

sequencing further showed that ML could predict ~80% of phenotypes within 2.5 hours, though performance varied by species. Notably, *Pseudomonas aeruginosa* showed limited ML accuracy due to dataset size and biological complexity underscoring the need for continued model expansion and refinement (1).

Conclusion

WGS-based genomic AST represents the most compelling path toward faster and more accurate susceptibility prediction, offering a clear advantage over both phenotypic and marker-based molecular methods. While phenotypic AST remains the clinical gold standard, its turnaround time is incompatible with the urgency of bloodstream infections. Likewise, marker-based prediction is valuable as supporting context but cannot serve as a standalone diagnostic, given its limited scope and inability to capture complex or emerging resistance mechanisms. In contrast, validated ML models trained on sufficiently large and diverse datasets can deliver accurate S/I/R predictions directly from early sequencing data, enabling clinicians to refine therapy in hours, rather than days, before conventional results are available.

As sequencing technologies become faster, portable, and increasingly integrated with automated workflows, including microbial DNA enrichment, real-time sequencing, taxonomic classification, and ML-driven prediction, the clinical feasibility of genomic AST continues to grow. Integrating these systems with hospital information systems and stewardship platforms will further enhance their impact. Looking ahead, expanding training datasets remains essential, particularly for regions facing a heavy burden of antimicrobial resistance such as India. Incorporating locally prevalent resistance mechanisms into global ML models will strengthen their robustness and equity. This underscores the importance of data-sharing partnerships and collaborative initiatives to acquire high-quality genomic and phenotypic datasets worldwide.

Bibliography

1. BacT-Seq, a Nanopore-Based Whole-Genome Sequencing Workflow Prototype for Rapid and Accurate Pathogen Identification and Resistance Prediction from Positive Blood Cultures: A Feasibility Study. Azami, Meriem El, et al. 1, 2026, *Diagnostics*, Vol. 16, p. 133.
2. Correlation between phenotypic antibiotic susceptibility and the resistome in *Pseudomonas aeruginosa*. Jaillard, Magali, et al. 2, 2017, *International journal of antimicrobial agents*, Vol. 50, pp. 210–218.
3. Evaluation of a machine learning system for genomic antimicrobial susceptibility determination on a clinically representative test set. Wittenbach, Jason D., et al. 2026, *Microbiology Spectrum*, pp. e00564–25.
4. Predicting bacterial resistance from whole-genome sequences using k-mers and stability selection. Mahé, Pierre and Tournoud, Maud. 1, 2018, *BMC bioinformatics*, Vol. 19, p. 383.

Use case 3

Sada Santé: A Franco–Indian Pilot for Privacy–Preserving AI in Cardiovascular Health

by iSPRIT Foundation.

Co-authored by:

Sharad Sharma, Co-founder, iSPRIT Foundation.

Harshit Kacholiya, Tech Policy and Strategy at iSPIRT

Sunu Engineer, Chief Technology Officer at Ryussi Technologies (P) Ltd.

Saumya Jetley, Faculty, Plaksha University Ph.D. in Computer Vision and Machine Learning Algorithms, University of Oxford

Sada Santé (“forever health”) is a flagship pilot under the France–India AI for Health initiative, jointly undertaken by iSPIRT India and the Health Data Hub (HDH) at Paris Santé Campus. The pilot reflects a shared ambition to advance responsible, large-scale medical artificial intelligence by aligning India’s strengths in population-scale data and digital public infrastructure with France’s leadership in clinical research, ethics-by-design, and AI governance.

The pilot focuses on cardiovascular disease, beginning with myocardial infarction (MI), a major cause of morbidity and mortality in both countries. By bringing together cardiology datasets from India and France, Sada Santé explores whether AI models trained across diverse populations and health systems can improve early risk prediction, stratification, and understanding of complex, multifactorial disease patterns. Beyond MI, the pilot is intentionally designed as a reusable framework for studying other complex health conditions that depend on longitudinal, multimodal data, including clinical records, imaging, biomarkers, and contextual factors.

A defining feature of Sada Santé is its privacy-preserving, governance-by-architecture approach to cross-border data collaboration. The pilot uses DEPA-enabled (Data Empowerment and Protection Architecture) data-sharing mechanisms to ensure that participating institutions retain control over their data at all times. Raw data does not move and is never directly accessed by external parties. Instead, data is combined and analyzed only within Confidential Clean Rooms (CCRs) operating under strict technical, legal, and organizational safeguards. Computation occurs within these secure environments, with strong differential privacy constraints applied throughout, ensuring that only trained models and approved aggregate outputs are made available to researchers and model developers.

Trust, accountability, and clinical relevance are central to the pilot’s design. Sada Santé incorporates end-to-end provenance tracing, enabling full auditability of data contributions, model training processes, and model versions. This supports a robust matrix of validation, certification, and reproducibility checks prior to any broader deployment. Explainability is treated as a core requirement, ensuring that resulting models can be meaningfully interpreted by clinicians, researchers, and regulators.

Strategically, Sada Santé serves as a proof-of-concept for scalable, privacy-preserving medical AI collaboration between France and India. It demonstrates how harmonized governance, rather than uniform regulation, can enable trusted cross-border innovation. As a pilot, it lays the groundwork for future Franco-Indian initiatives in AI for health that are scientifically rigorous, ethically grounded, and globally relevant.

5. Recommendations

The following recommendations are derived from the strategic priorities, observations, and use cases above. They focus on policy-level actions and institutional mechanisms that can be jointly advanced by India and France. The examples and actions provided are indicative and non-exhaustive, and are intended to illustrate feasible pathways for implementation rather than prescribe technical solutions. These recommendations can be operationalized through alignment between existing national, European, and bilateral funding instruments, including Horizon Europe, ANR, DST, ICMR, and relevant Indo-French cooperation programs.

Establish a Shared Indo-French Governance and Ethics Framework for Medical AI

Rationale

Medical AI systems currently lack a harmonized governance and validation framework comparable to those used for pharmaceuticals or medical devices. Divergent ethical, consent, and accountability approaches create uncertainty for developers, clinicians, regulators, and patients, limiting trust and responsible deployment.

Mechanism

1. Establish a joint Indo-French ethics and governance mechanism for medical AI collaborations.
2. Align existing national, European, and international ethical frameworks rather than creating new regulatory instruments.
3. Define shared principles for accountability, oversight, and post-deployment monitoring of medical AI systems.

Policy-level solution

A coordinated governance framework that provides clarity, consistency, and trust across the lifecycle of medical AI used in healthcare settings.

Proposed Indo-French Deliverable (12–36 months)

An Indo-French Medical AI Governance and Validation Framework, developed in coordination with relevant French and Indian health and digital authorities, operationalized through a permanent Indo-French Medical AI Governance Board. This framework would be piloted initially in brain health and tertiary hospital networks (e.g. AP-HP, Paris Brain Institute, AIIMS, and IIT-affiliated hospitals) before being extended to other clinical domains.

Enable Privacy-Preserving Data Collaboration Through Federated Approaches

Rationale

Strict data protection and privacy regulations, while essential, currently constrain cross-border research and limit access to diverse and representative health data. This reduces the generalizability and robustness of medical AI systems intended for clinical use.

Mechanism

1. Jointly advance federated learning approaches that enable collaborative model development without direct data transfer.
2. Align data standards and interoperability frameworks compatible with General Data Protection Regulation (GDPR) and the Digital Personal Data Protection (DPDP) Act.

3. Treat federated approaches as shared research infrastructure rather than isolated pilot projects.

Policy-level solution

A trusted, privacy-preserving collaboration model that enables cross-border research while respecting existing legal and ethical constraints.

Proposed Indo-French Deliverable (12–36 months)

An Indo-French Federated Health Data & AI Research Infrastructure, jointly governed and interoperable with both GDPR and India's DPDP Act. This infrastructure would initially support brain health and neurological datasets and then expand to other disease areas.

Define Standardized Translational and Validation Pathways for Medical AI

Rationale

Many AI systems demonstrate promise in research environments but fail to progress into routine clinical use due to the absence of standardized validation and translational pipelines.

Mechanism

1. Jointly define stages for progression from methodological research to clinical evaluation and deployment.
2. Support multicenter testing and population-level validation across diverse care settings.
3. Encourage post-deployment monitoring to assess real-world performance and safety.

Policy-level solution

Clear and predictable validation pathways that reduce uncertainty and improve the credibility and uptake of medical AI in health systems.

Proposed Indo-French Deliverable (12–36 months)

An Indo-French Clinical AI Validation Network linking AP-HP, AIIMS, and partner hospitals, using shared protocols and federated evaluation pipelines. This network would be piloted in brain imaging, neurology, and neuro-rehabilitation before broader deployment.

Strengthen Human Capacity, Training, and Trust Within Health Systems

Rationale

The effective adoption of medical AI depends on the confidence, preparedness, and trust of clinicians and health system users. Technical performance alone is insufficient without adequate human capacity.

Mechanism

1. Strengthen joint training modules focused on usability, explainability, and limitations of medical AI systems.
2. Support short courses, exchange programs, and mobility initiatives for clinicians and researchers.
3. Promote user-centered design principles aligned with clinical workflows.

Policy-level solution

Health systems equipped with the skills and confidence required to responsibly integrate AI into clinical and public health practice.

Proposed Indo-French Deliverable (12–36 months)

An Indo-French Medical AI Academy, embedded within universities and teaching hospitals. This academy would ensure that AI adoption is driven by clinical understanding and trust, not just technical availability.

Embed Equity and Generalizability as Core Criteria for Deployment

Rationale

AI systems trained on narrow or homogeneous datasets risk reinforcing existing disparities in access to care and health outcomes.

Mechanism

1. Promote inclusion of diverse populations, languages, and care settings in data generation and validation.
2. Encourage population-level evaluation and continuous performance monitoring.
3. Integrate equity considerations into deployment and oversight criteria.

Policy-level solution

AI-enabled healthcare systems that perform reliably and equitably across populations rather than amplifying existing inequalities.

Proposed Indo-French Deliverable (12–36 months)

An Indo-French Population-Scale AI Equity & Performance Observatory, designed to ensure fairness and real-world reliability. This observatory would position Indo-French cooperation as a leading example of responsible, population-level AI in health.

Build Sustained Scientific, Educational, and Innovation Ecosystems

Rationale

Current collaborations are often project-based and time-limited, limiting continuity, knowledge transfer, and long-term impact.

Mechanism

1. Establish joint graduate, doctoral, and postdoctoral programs between Indian and French institutions
2. Support long-term researcher mobility and institutional partnerships
3. Coordinate funding mechanisms and shared validation platforms

Policy-level solution

Durable Indo-French ecosystems that link education, research, validation, and innovation to ensure sustained impact over time.

Proposed Indo-French Deliverable (12–36 months)

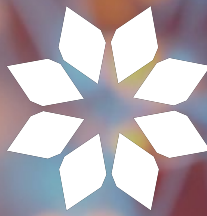
An Indo-French AI for Health Flagship Program, jointly funded and institutionally anchored, including:

1. Co-funded PhD, MD-PhD, and postdoctoral tracks in medical AI and digital health,
2. Permanent Indo-French research hubs linking Sorbonne University, Paris Brain Institute, Inserm, AIIMS, and IIT networks,
3. Shared clinical innovation and validation platforms,
4. And aligned funding calls between Horizon Europe, ANR, DST, and ICMR.

This flagship would ensure that Indo-French cooperation in medical AI becomes structural, cumulative, and globally visible, rather than remaining project-based.

References

- World Health Organization. Ethics and Governance of Artificial Intelligence for Health. Geneva; 2021.
- World Health Organization. Global Strategy on Digital Health 2020–2025. Geneva; 2020.
- Organisation for Economic Co-operation and Development. Artificial Intelligence in Society. Paris; 2019.
- Organisation for Economic Co-operation and Development. Health Data Governance: Privacy, Monitoring and Research. Paris; 2015.
- European Commission. Regulation of the European Parliament and of the Council Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act). Brussels; 2024.
- European Commission. European Health Data Space. Brussels; 2022.
- NITI Aayog. National Strategy for Artificial Intelligence: #AIforAll. New Delhi; 2018.
- Government of India. Digital Personal Data Protection Act. New Delhi; 2023.
- World Health Organization. WHO Guidance on Regulatory Considerations on Artificial Intelligence for Health. Geneva; 2023.
- The Lancet Digital Health Commission. Governing Health Futures 2030: Growing Up in a Digital World. London; 2021.



FRANCE INDIA
FOUNDATION

AI & Automotive

France India AI Initiative
February 2026

AI & Automotive

Heads of AI & Automotive working group



Sarita Kaloya

Senior Director and Global Databricks & Azure
Practice Leader at Capgemini,
Young Leader.



Neha Arolkar

Director India Client Partner & Global Innovation
Leader, Automotive at Capgemini,
Young Leader.

1. Strategic Review and Priorities

The automotive industry is no longer defined by the combustion engine, but by the code that governs it. The ongoing transformation in the industry is the harbinger of a future largely driven by artificial intelligence (AI), especially in autonomous driving, vehicle safety, manufacturing optimization, and personalized user experiences. This paper explores the strategic role of AI in Automotive, with a specific focus on the collaboration between India and France.

As we navigate the Software-Defined Mobility era, the geopolitical synergy between France and India offers a unique competitive advantage. Under the Horizon 2047 roadmap, our nations have committed to a partnership rooted in security, sovereignty, and sustainability¹.

This initiative aims to go beyond adopting Artificial Intelligence into establishing AI Sovereignty in mobility. Currently, the global automotive AI narrative is dominated by the US (Silicon Valley software models) and China (battery dominance and aggressive deployment of electric & autonomous vehicles). France and India, standing together, represent a third pole:

- France brings a century of precision engineering, luxury heritage, and European regulatory maturity (GDPR/EU AI Act).

¹ Horizon 2047: 25th Anniversary Roadmap [Online] / auth. (France) Ministry of External Affairs (India) & MFA. - 2023-2025. - <https://www.mea.gov.in/bilateral-documents.htm?dtl/36806/>

[Horizon_2047_25th_Anniversary_of_the_IndiaFrance_Strategic_Partnership_Towards_A_Century_of_IndiaFrance_Relations](#)



- India brings the world's largest pool of digital talent, a massive data generation ecosystem (3rd largest auto market), and a culture of frugal innovation. Recently, Stanford University ranked India 3rd on its AI Vibrancy Index 2024, just behind the US and China².

The Indo-French AI collaboration presents a unique opportunity to leverage both countries' strengths in technology innovation, manufacturing capabilities, and engineering expertise. The AI for Automotive Accelerator offers a unified space to create & showcase breakthrough innovations, influence policy and collaborate across borders on the shared goals of:

- Safe, sustainable mobility
- Infrastructure connectivity
- Workforce readiness

Under its ambit, this first white paper covers key trends, opportunities, use cases, actionable insights and mutually beneficial policy-oriented recommendations for a collaborative path forward on AI strategy in the automotive sector.

2. Scope

For an end-to-end transformation, our scope needs to encompass distinct types of key players.

2.1. Industry

- **Automotive Original Equipment Manufacturers (OEMs) and their dealers:** going beyond building cars to shaping the mobility experience (Stellantis, Renault, Mahindra)
- **Tier 1 & 2 suppliers:** critical layer that manufactures smart components (Michelin, Valeo)
- **Mobility-as-a-Service providers:** Fleet managers and rentals (Arval, Europcar, Uber, Ola), multi-modal and last mile integrated mobility providers
- **Tech Partners & System Integrators:** The bridge builders (e.g., Capgemini, Dassault Systèmes) facilitating the IT-OT convergence
- **Deep-Tech Startups:** Agile players in Autonomous Vehicles, battery analytics and computer vision

2.2. Academia & Research Hubs

The IITs (India) and INRIA/CNRS/Grandes Écoles (France) joint research corridors, incubators such as NSRCEL – IIM Bangalore

2.3. Government Bodies

Ministry of Heavy Industries (India), Ministry of Economics/PFA (France) for regulatory harmonization, Skill Development Corporations (e.g. TNSDC).

Impact of AI in Automotive spans 4 interconnected domains

(with non-exhaustive examples)

² Human-centred Artificial Intelligence: Global AI Vibrancy Index [Online] / auth. University Stanford. – 2025. – <https://hai.stanford.edu/ai-index/global-vibrancy-tool>.

- Vehicle Intelligence (Autonomous Driving & ADAS (Advanced Driver Assistance Systems))
- Operations (Manufacturing Optimization, Predictive Maintenance)
- Customer Experience (In-cabin and buyer/user journey personalization)
- Mobility Services Ecosystem Orchestration (Fleets, rental, dealer platforms)

Core Functional Domains of AI in Automotive

- **Autonomous Vehicles (AVs) and ADAS:** AI enables perception, motion planning and control for autonomous functions and ADAS (real-time object detection, lane-keeping, adaptive cruise control and accident avoidance)
- **Connected Platforms:** Telemetry-driven analytics, predictive maintenance, remote diagnostics, over-the-air (OTA) updates, and connected infotainment systems extending into usage-based insurance and service recommendations
- **Smart Manufacturing & Service Operations:** AI-optimized manufacturing and aftermarket processes – robotics, quality control, energy optimization, parts planning, turnaround time etc.
- **EV Ecosystems:** AI-optimized charging availability, EV range, improving battery management, managing energy efficiency at fleet-level to optimize cost
- **Mobility Services Intelligence:** AI-powered demand forecasting, intelligent routing, utilization, dynamic pricing & valuation, automated damage detection & claims, personalized vehicle matching, multi-modal & last-mile journey planning and much more.

AI Adoption Drivers in the Automotive Sector

- **Safety, Cybersecurity, Compliance:** meeting evolving mandates
- **Cost & Uptime Management:** lower Total Cost of Ownership (TCO), faster turnaround
- **Platform and Economics:** new revenue streams from software-defined features
- **Experience and Convenience:** personalization across purchase, usage and services
- **Sustainability:** optimized routes, charging, battery life

The deployment of AI in the automotive industry is taking place both at the edge (on the vehicle, for real-time safety & control) and in the cloud (for learning loops and orchestration). Google Maps is a well-known historical example.

3. Observations

3.1. Market Trends

Current State of AI across the Automotive Value Chain

Parameter	France (Precision Hub)	India (Scale & Software Hub)
Primary Focus	Industrial Metaverse, Industry 4.0, Sovereign Cloud	Connected Services, Fleet Analytics, Affordable EV Tech
AI Maturity	High in Manufacturing (e.g. Predictive Maintenance) and R&D	High in Consumer Apps, Infotainment & Aftersales
Regulation	Strong Framework – EU AI Act. Focus on data privacy, liability and ethics	Evolving Framework – DPDP Act. Focus on ease of innovation vs privacy
Talent Pool	Deep domain experts	Massive volume of data and AI/ML engineers
Consumer Demographics	Ageing buyers	Young, tech-savvy

3.2. Expert Points of View

Yves Caseau³, former CIO and current Head of Digital & Information Systems at Michelin Group, states that Michelin is using AI for optimizing tyre manufacturing. They are also exploring designs – such as that of the new winter tyre Alpin 7 – through algorithms, AI simulations and digital twins. However, he cautions against relying blindly on AI and encourages using AI tools to progress quickly while retaining our human learning and problem-solving capabilities.

Jean-Marie Lapeyre⁴, Chief Technology & Innovation Officer, Automotive Industry, Capgemini opines that Large Language Models (LLMs) are increasingly in the spotlight due to enormous funding, however they are a small piece of the large AI world. LLMs are designed to probabilistically anticipate, not formally reason. The automotive industry needs formal verification and validation in ADAS from a safety standpoint, therefore LLMs themselves should not be entrusted with driving cars. Enter Perception Models – these can process large amounts of information, identify patterns and distil insights, often better than humans. Renault is a good example of how both LLMs and agents are used to address different use-cases for the car.

Dr. Luc Julia⁵, Chief Scientific Officer at Renault Group and co-creator of Siri, explains that the French automaker uses AI primarily in two domains: outside the car (safety, autonomous driving, vehicle-to-X connectivity) and inside the car (voice assistants, driver alertness monitoring, etc.). Renault uses

³ Conversations for Tomorrow (90-103) [Online] / auth. Institute Capgemini Research. – December 2025. – <https://www.capgemini.com/wp-content/uploads/2025/11/Final-Web-Version-Report-CFT10.pdf>

⁴ [Interview] / interv. Lapeyre Jean-Marie. – 2025.

⁵ Dr Luc Julia [Online] / auth. Capgemini Research Institute. – <https://www.capgemini.com/insights/research-library/a-conversation-with-dr-luc-julia/>.

Generative Artificial Intelligence (GenAI) to offer a human-like user interface & experience, while leveraging agentic AI on the system architecture side. According to Dr. Julia, the OEM started working on a GenAI-powered intelligent assistant for the recently released Renault 5 electric model even before ChatGPT became ubiquitous in 2023. On the rising topic of agentic AI, he thinks of the car as an ecosystem of specialized, semi-independent, efficient AI agents with a master orchestrator to decide the overall actions.

3.3. Current Challenges and Future Opportunities for France and India regarding AI in the Automotive sector

Despite the potential, three key friction points exist.

3.3.1. Data Sovereignty, Privacy and Security Concerns

Autonomous vehicles and connected car systems generate massive amounts of data, raising concerns over data privacy and cybersecurity.

French OEMs are hesitant to store R&D data on non-sovereign clouds; Indian regulations are increasingly mandating local data residency. We lack a "Trusted Data Corridor."

3.3.2. Legacy Debt and Hardware Limitations

French factories have decades of legacy systems (PLCs/SCADA) that are hard to modernize; Indian OEMs are newer but struggle with high capital costs for Industry 4.0 hardware. Edge devices on vehicles face constraints in processing power, memory and energy consumption. This necessitates the development of specialized AI chips that balance performance with efficiency.

3.3.3. The PoC Purgatory

Both nations suffer from AI initiatives getting stuck in Proof-of-Concept stages i.e. unable to scale from pilot to production. There are various reasons, such as a lack of clear ROI definitions in the manufacturing sector, solutions that work in the lab but fail in the real world, and investor unwillingness to fund for longer periods. Automotive AI solutions today are highly fragmented, leading to difficulties in standardizing AI models across vehicles and platforms. In general, the European automotive industry has developed fewer standards compared to other industries, leading to comparatively higher costs for non-differentiated components.

However, these same challenges present opportunities for alignment and technological synergy.

Challenge	Opportunity	Franco-Indian Collaboration	Risks mitigation
Data Sovereignty	Sovereign Federated Learning Alliance	Regulatory gold standards from France, diverse datasets from India with data-savvy workforce for annotation within a Federated Learning Network. Model weights shared back.	To prevent catastrophic 'forgetting' of basic rules, maintain separate fine-tuned layers for different geographies within the base model
Legacy Debt and Capex	Brownfield AI Modernization	Advanced hardware capabilities and design expertise from France, "frugal AI" retrofittable kits from Indian startups tailored to international industrial specifications.	To minimize interoperability issues between modern AI edge devices and legacy machines, use standardized middleware bridge
Scaling beyond PoC	Joint Validation	World-class simulation and digital twin capabilities from France, real-world stress-testing environments in India. Joint definitions and certifications for quick exit of successful solutions from the Industrial Metaverse for further scaling.	Designate and make testing zones 'IP-safe' with laws of origin country applicable, to prevent leakage of proprietary algorithms

The anticipated impact ranges from improved edge-case handling while avoiding expensive physical infrastructure, to extreme digital and physical testing for safety at true scale.

4. Use cases

In accordance with our objectives, we have identified 3 priority use cases to explore as pilots or joint workstreams.

⁶ EU Road Safety: Towards Vision Zero [Online]. - 2022. - https://cinea.ec.europa.eu/publications/digital-publications/eu-road-safety-towards-EV-vision-zero_en.

⁷ [Online]. - Business Standard, 2024. - https://www.business-standard.com/india-news/morth-aims-to-halve-road-accidents-and-casualties-by-2030-gadkari-123121100964_1.html.

⁸ Charging Infrastructure in India 2025: Bold Growth & Gaps (Tata Power & Statiq Operator Disclosures. (Q1 2025)) [Online] // greenglobe25.in. - 2025. - <https://greenglobe25.in/ev-charging-infrastructure-in-india-2025/>.

Objective	Use Case and Details	Impact Metrics
Safe, sustainable mobility	Safety in Autonomous Driving - Object Detection and Classification: Using AI to recognize pedestrians, cyclists, other vehicles, and road obstacles. - Lane Keeping & Collision Avoidance: Advanced AI systems can detect lane boundaries and avoid collisions by steering or braking autonomously. - Predictive Maintenance for AVs: AI models can predict vehicle component failures before they occur, improving safety and reducing maintenance costs. - Driver Behavior Analytics: AI systems analyze user driving patterns to deliver customized suggestions for efficiency and safety.	- Help achieve 50% reduction in Road Fatality Rate by 2030 in line with Vision Zero⁶ (EU) and MoRTH⁷ (India) - Increase ADAS Corrective Action Accuracy to reduce safety feature turnoff rate by drivers
Improved EV Connectivity & Infrastructure	EV Connectivity and Battery Optimization - Charging Infrastructure: AI-driven predictive solutions for Charge Point Operators (CPOs) to increase availability, compatibility and reliability of charging equipment with a unified app interface. - Battery Management Systems (BMS): AI-powered energy optimization algorithms to improve range and efficiency of EVs, provide contextualized range prediction, battery health score for resale value and enable battery lifecycle transparency supporting recent Battery Passport and BPAN initiatives. - Vehicle to Grid (V2G) integration in France: V2G balancing yield (electricity credits), treating EV batteries as Mobile Storage Units in the national energy trading market.	- Increase Charger Uptime Reliability from current 68% average in India⁸ - Enable up to 10% peak-load reduction on grid in France
Workforce transformation	Talent Pipeline Readiness - Cross-skilling for students: curriculum upgrades to cover both technical, domain and AI topics in depth, along with applications at the intersection of these themes. - Up-skilling for entry-level workforce: basic AI literacy and applied AI combined with existing technical or functional expertise. - Re-skilling for senior leaders: programs designed for experienced professionals to unlearn and re-learn new-age AI skills to stay ahead of the curve.	% of non-Computer Science engineering/science graduates who have at least 20% AI credits in coursework % of "Day-1 AI-ready" talent by reduction in curriculum-to-industry lead-time

5. Recommendations

To operationalize the "Horizon 2047" vision, we propose the following recommendations, especially with a policy lens. These are designed to be high-impact, politically viable, and technically feasible, leveraging the complementary strengths of the French and Indian AI ecosystem.

5.1. Government (Policy & Infrastructure)

Launch a dedicated Franco-Indian Automotive AI Mission

Proposal

Create a joint mission focused on Edge AI and SDVs, Autonomous and assisted driving, and Automotive semiconductors and AI software platforms. Promote AI-driven policy frameworks.

Mechanism

Align with Digital India, Make in India, PFA (French Automotive Platform) and EU mobility strategies. Both governments could incentivize AI research, deployment, and adoption by providing grants, subsidies, and tax breaks for automotive companies that invest in AI technologies.

Impact

Long-term strategic leadership in automotive AI.

Shared Indo-French Mobility Data Corridor and Infrastructure Access

› Project "Setu-Pont"

Proposal

Establish a legally compliant, encrypted data pipeline that allows French OEMs to train algorithms on Indian road data (for edge-case learning) and Indian startups to access French synthetic data (for validation) without compromising data sovereignty.

Mechanism

Utilize Federated Learning: Instead of moving raw data across borders (violating GDPR or DPDP), the model travels to the local data center, learns, and returns only the updated weights. Create a White-Listed Data Zone: A regulatory sandbox overseen by CNIL (France) and the Data Protection Board (India) where specific automotive R&D data can be exchanged under a simplified compliance regime. Also develop cross-border standards for AI-driven data governance to ensure privacy protection and cybersecurity in connected and autonomous vehicles.

Impact

French AVs learn to handle chaotic traffic in India to become robust; Indian EVs get access to mature battery longevity data from France.

› Sharing AI infrastructure with a Sovereign Compute Access Pass

Proposal

A quota-based system allowing certified joint-venture startups to access high-performance computing (HPC) resources in both nations.

Mechanism

India has approved the IndiaAI Mission (₹10,372 Cr)⁹ to build GPU infrastructure. France has the Jean Zay supercomputer and regional AI clusters. Allocate 10% of reserved compute time for Indian and French strategic startups working on decarbonization, safety, etc. Establish national automotive AI testbeds – Indian driving datasets, Simulation and validation platforms.

Impact

Lowers the barrier to entry, innovation & commercialization for deep-tech mobility startups that currently spend significant capital on cloud/GPU costs from US providers.

› "Sadak-Rue": The Open Edge-Case Dataset

Proposal

Creation of the world's largest open-source dataset for chaotic traffic scenarios, hosted on the IndiaAI Datasets Platform AIKosh.

Mechanism

Data collection of "edge cases" (cows on roads, wrong-way drivers, heavy rains, narrow European cobblestone streets). Anonymized and tagged by Indian data annotation firms, validated by French research institutes.

Impact

Breaks the monopoly of US/Chinese tech giants on AV training data. It positions the Indo-French bloc as a global hub for robust AI training.

Mutual Recognition of AI Safety Certification with Tax Incentives for Safety & Cybersecurity

Proposal

Firstly, an agreement between UTAC (testing & certification body in France) and ARAI (Automotive Research Association of India) to mutually recognize AI-based safety tests. Secondly, in India we propose a lower GST bracket (e.g. 5% instead of 28%) for vehicles that score above a certain threshold on AI-Safety Reliability (as certified by the mutual ARAI-UTAC framework).

Mechanism

For example, if a Driver Monitoring System (DMS) is certified compliant with the EU General Safety Regulation (GSR) by UTAC, it should be fast-tracked for AIS (Automotive Industry Standards) approval in India, and vice-versa. Aim to harmonize regulations to the extent possible. Joint development of a new "Turing Test for Cars": A standardized test track protocol (physical and virtual) to evaluate how AI handles ethical dilemmas (e.g. unavoidable accident scenarios).

Impact

Can potentially reduce time-to-market for new models by 12-18 months while eliminating redundant testing costs for exporters and reducing prices for consumers. Faster innovation with global compliance.

⁹ Cabinet Approval: IndiaAI Mission (₹10,372 Cr) [Online] / auth. Press Information Bureau (PIB) Govt. of India. – 2024. – <https://www.pib.gov.in/PressReleaselframePage.aspx?PRID=2012355®=3&lang=2>.

Incentivize Edge AI and Semiconductor Development

Proposal

Extend PLI and fiscal incentives to automotive AI chips, NPUs, and edge computing platforms. Promote Indo–French joint ventures in chip design, fabrication, and packaging to build a secure semiconductor ecosystem.

Mechanism

Align with India's Semiconductor Mission and France's semiconductor strategy. Create a PLI sub-category for automotive-grade AI chips with safety and efficiency benchmarks. Set up joint OSAT facilities and design labs for low-power AI inference chips. Enable public–private partnerships for R&D and tax credits. Establish a "Sovereign Silicon" framework with joint certification standards for compliance.

Impact

Strengthens technology sovereignty and reduces reliance on external ecosystems. Builds a resilient supply chain against global disruptions. Powers real-time edge AI for autonomous driving with functional safety. Positions Indo–French collaboration as a global hub for automotive AI hardware innovation.

5.2. Industry (OEMs, Suppliers, Tech Companies)

Leverage GCCs as Innovation Hubs

Proposal

Shift from the traditional outsourcing model to a co-creation model. This is already happening to some extent as major French OEMs such as Renault and Stellantis are upgrading their Indian Global Capability Centers (GCCs) from back-offices to hubs of innovation as an extension of their headquarters in France.

Mechanism

Real-time Digital Twins: Consider a hypothetical physical EV battery lab in Grenoble that has a digital twin in Pune. Experiments run physically in France generate data that is analyzed in real-time by AI agents in India, which then suggest parameter changes for the next physical run in France. Shared IP ownership for innovations development in these labs.

Impact

Accelerates R&D cycles by enabling "follow-the-sun" ways of working i.e. across time-zones, with seamless hand-offs.

Standardization of the Industrial Metaverse

Proposal

Adoption of a common data standard for factory interoperability, similar to the Catena-X initiative but tailored for the Indo–French ecosystem.

Mechanism

Agreeing on open protocols (e.g., Asset Administration Shells) so that a machine in a Chennai factory can "talk" to a supply chain algorithm in Paris without complex custom integration. Joint working groups

- for example, between Alliance Industrie du Futur (France) and SAMARTH (Smart Advanced Manufacturing And Rapid Transformation Hub) Udyog (India).

Impact

Enables true supply chain transparency and carbon footprint tracking (Scope 3 emissions) for export compliance.

Strengthening the Supply Chain with and for AI

> The SME AI Uplift (Tier 2/3 Modernization) and Co-creation among OEMs & Startups

Proposal

A private-sector fund, co-anchored by Tier-1 giants (e.g. Valeo), to subsidize AI adoption for smaller Tier 2/3 suppliers (e.g. in Hauts-de-France cluster, Motherson Sumi in India)

Mechanism

"AI-in-a-Box" Kits: Deploying simple pre-trained computer vision cameras on the assembly lines of small component manufacturers making gears, fasteners etc. to detect defects. Subsidized by the AI Booster (France 2030)¹⁰ and Indian PLI schemes if the SME is part of the Indo-French supply chain. Establish co-development programs where start-ups work directly with OEMs and Tier-1 suppliers from early design stages. Encourage long-term partnerships over pilots.

Impact

Reduces the risk of quality deficit. A 2% defect rate at a Tier-3 supplier can shut down a Tier-1 line, however AI-based visual inspection could potentially cut this to near zero. Overall reduced time-to-market and better integration of innovation into production vehicles.

> 'Mobile Battery' Trading License

Proposal

Recently, 65% of French EV owners have expressed interest in purchasing a V2G-capable vehicle¹¹. A new class of 'Micro-Utility' licenses, designed by and for EV fleet owners in France, could benefit them.

Mechanism

In alignment with the French Ministry of Economics and Finance, and the French Energy and Regulatory Commission (CRE), fleet owners can design and leverage such a license to trade electricity back to the grid via AI agents.

Impact

Alternative revenue stream for public transport operators.

¹⁰ Launch of the AI Booster France 2030 program [Online] / auth. France Government of. - 2023. - <https://www.entreprises.gouv.fr/la-dge/actualites/lancement-du-programme-ia-booster-france-2030>.

¹¹ Consumer Monitor 2023 Country Report: France [Online] / auth. Commission European Alternative Fuels Observatory (EAFO) / European. - 2023. - <https://alternative-fuels-observatory.ec.europa.eu/system/files/documents/2024-07/FR%20Report%20-EC%20LAYOUT-FINAL.pdf>.



› Semiconductor supply resilience

Proposal

Semiconductor companies, cloud providers, and system integrators should collaborate to deliver AI-optimized automotive SoCs and integrated edge-cloud AI stacks.

Mechanism

Encourage India–France industrial collaboration in semiconductor design, testing, and validation. Leverage trusted & shared data (see Data Corridor Project “Setu-Pont” recommendations in previous section) to improve forecasting and mitigate risks of supply shocks. Public–Private Partnerships can ease the financial burden of innovation.

Impact

Reduced dependency on external supply chains and improved resilience.

5.3. Academia & Research

Establish Indo–French Center of Excellence in Automotive AI

Proposal

A consortium of currently active French and India universities, research institutes and on-campus incubators to enhance our collective AI R&D capabilities.

Mechanism

Focus on driving research & development on next-generation AI solutions for vehicles – covering autonomous systems, Edge AI optimization, functional safety and AI ethics. Co-funded by governments and industry, backed by the Franco–Indian Automotive AI Mission. Position India and France as global leaders in AI and mobility by collaborating on international AI policy discussions, establishing joint patents, and creating new AI standards for vehicle safety and performance.

Impact

World-class applied research ecosystem enabling both nations to stay ahead of the curve.

Large–Scale Skilling Programs

› Student Up–skilling and Cross–skilling Programs

Proposal

Equipping engineering or technical undergraduate students with real-world skills, combining technical/functional knowledge with basic AI literacy. This can be followed by specialization in various aspects of AI with a focus on applying AI in real-world business contexts.

Mechanism

Update engineering and technical undergraduate programs’ curriculum with specialized modules in Embedded AI and real-time systems, Automotive software engineering and Autonomous driving algorithms. Encourage Indo–French university exchange programs, students’ active participation in industry-sponsored hackathons & government-provided courses and building their own portfolio of AI-based side-projects.

Impact

Day-1 Industry-ready talent, able to be deployed on live projects for quick results.

› **Workforce Reskilling Programs**

Proposal

National skilling initiatives and industry/company-specific reskilling programs.

Mechanism

As working engineers on combustion-engine vehicles transition to software-defined electric vehicles powered by AI, they need to be rapidly upskilled before becoming obsolete. Traditional developers and technical leads need to also function as automotive software and safety professionals.

Impact

Rapid scaling of skilled workforce.

› **Executive AI Ethics & Management Exchange**

Proposal

A rotation program for senior technical and business leaders.

Mechanism

French leaders spend 3 months in India across tech companies shadowing counterparts leading scaled teams; vice-versa for Indian leaders with an immersion in the French research hubs. Curriculum focuses on the cultural nuances of AI deployment: GDPR vs. DPDP, Luxury vs. Value engineering.

Impact

Prevents cultural friction in joint ventures and ensures leadership understands the "Sovereign AI" imperative.

Joint Internships and Fellowships linking student talent pipeline with industry

Proposal

Up to 100 dedicated fellowships, internships, PhD and post-doc programs funded by jointly by CEFIPRA (Indo-French Centre for the Promotion of Advanced Research) and by Industry (OEMs, suppliers etc.) focused specifically on frugal AI for mobility. Encourage research aligned with industry use cases and commercialization.

Mechanism

Research labs pick a relevant problem statement e.g. how to run advanced pedestrian detection algorithms on low-cost, low-power chips which is crucial for affordable cars and 2-wheelers in India. Students work on addressing these problems and get opportunities to implement their solutions directly with the industry players or to 'spin off' their solution as a startup in associated incubator hubs.

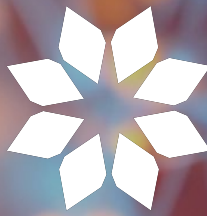
Impact

Creates a pipeline of "AI-ready" talent who understand and have firsthand experience in both the automotive domains as well as AI.

Conclusion

In a shifting global landscape, the symbolic synergy of the Gallic Rooster and the Indian Peacock stands tall as an anchor of sovereignty. Horizon 2047 is a commitment to lead together. Transportation & mobility along with technological innovation are key components of the strategic dialogue between France and India.

We believe that while Artificial Intelligence is not a silver bullet for all our challenges, it is not a temporary technology hype either. AI is a critical accelerator for both our nations to own our technology as a form of intelligent and resilient interdependence. Together, we can drive towards a brighter shared future. .



FRANCE INDIA
FOUNDATION

AI & Regulation

France India AI Initiative
February 2026

AI & Regulation

Heads of AI & Regulation working group



Ahmed Baladi

Partner, Gibson Dunn,
Young Leader.



Laurie-Anne Ancenys

Partner, Head of Tech & Data
Practice, Paris Office, A&O
Shearman,
Young Leader.



Chinmay Tumbe

Faculty Member, Economics
Area, Indian Institute of
Management, Ahmedabad,
Young Leader.

1. Strategic Review and Priorities

Artificial Intelligence (“AI”) has evolved beyond a mere technological breakup into a strategic asset that underpins economic growth, global influence and national security. Its applications now span across every major sector (from healthcare, automotive, and finance to agriculture) and are increasingly embedded in individuals’ daily lives.

Regulation of AI does more than define norms, it directly shapes the innovation landscape, steers business strategies, and influences global competitiveness which ultimately determines where talent and firms choose to build. AI is increasingly taking shape as a contest between global powers (China, US, India, Europe), each seeking not only technological dominance but the power to establish the standard-setting AI model. Beyond this technological competition there are distinct regulatory approaches followed by global powers that reflect national/regional strategies to create the most efficient ecosystem for the development of AI. As France and India deepen cooperation on AI, it is essential to situate that partnership within this global regulatory landscape and to understand how different approaches can either catalyze or constrain the emergence of world-class AI ecosystems.

The EU Artificial Intelligence Act¹² (“EU AI Act”) is widely regarded as the world’s first comprehensive, horizontal regulatory framework for artificial intelligence. Conceived by EU lawmakers as a benchmark

¹² Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act)

for global AI governance, it aims to shape regulatory approaches both inside and outside the European Union.

While praised for its ambition and emphasis on protecting fundamental rights, user safety, and transparency, the EU AI Act faces sustained criticism centered on complexity, administrative burden, and the risk of stifling innovation and competitiveness with entrepreneurial energy shifting toward jurisdictions whose governance is more proportionate and adaptive. Insights from the other AI regulations, in India but also the United Kingdom and Singapore, all of which have pursued more adaptive, principles-based regulatory models, offer valuable insights. These examples show how regulatory environments that balance innovation and accountability can attract AI development and international collaboration.

In response, EU institutions are considering targeted revisions to simplify obligations and reduce compliance costs. These efforts fit within a broader streamlining agenda, notably the Digital Omnibus Regulation proposal, the Mario Draghi Report¹³, which seeks to reduce burdens, eliminate duplications, and support innovation and competitiveness.

It is also worth mentioning the strategic role of EU Guidance in enabling a practical and innovation friendly AI regulatory framework. Across the European Union, the regulation of artificial intelligence is evolving rapidly. A defining feature of the EU's approach is the European Commission's commitment to providing non-binding but authoritative guidance to support consistent and practical implementation of the regulation. The Commission's publication of Guidelines on the definition of an AI system exemplifies this effort. These guidelines are explicitly intended to help providers and other stakeholders determine whether a software system qualifies as an AI system under the EU AI Act, thereby ensuring that the rules can be applied clearly and effectively in practice. They are designed to evolve over time in response to new technologies, questions, and real world use cases, reflecting a flexible and pragmatic regulatory strategy. The Commission's willingness to publish interpretative materials, support stakeholder understanding, and update guidance as technologies advance signals of a regulatory posture focused not only on compliance, but on enabling innovation within a transparent and predictable framework. They are intentionally designed to evolve over time, adapting to new technologies, emerging questions, and real world use cases, thereby supporting a more coherent and workable implementation across the Union. Through this combination of non-binding rules and adaptive guidance, the EU positions itself as a global standard setter while actively supporting developers, businesses, and public bodies in navigating AI obligations. This approach underscores the Union's broader objective: ensuring that AI contributes positively to society and the economy, while making regulatory compliance both achievable and aligned with technological progress.

However, despite this constructive and proactive approach, the Commission's guidance cannot, by nature, alter the underlying complexity or prescriptiveness of the legislative framework. Unlike more agile models in jurisdictions such as the UK or Singapore, which rely on principle-based oversight and regulatory sandboxing, EU guidance helps interpret the rules but does not reduce compliance burdens nor increase regulatory flexibility yet. In that sense, they enhance clarity and usability but stop short of transforming the framework into one that is more innovation-conducive by design. This highlights a

¹³ Mario Draghi, [The Future of European Competitiveness](#), September 2024

structural gap: the EU's willingness to support innovation through guidance is evident, yet these measures may not be sufficient to fully counterbalance the regulatory rigidity that can constrain Europe's competitiveness in the global AI landscape.

This section examines why the EU AI Act, while pioneering, may not optimally foster innovation and how targeted refinements drawing from India and other international AI regulations such as in the UK and Singapore, could materially improve Europe's competitiveness.

1.1. AI Regulation in the European Union (EU AI Act)

1.1.1. Prescriptive obligations

The EU AI Act imposes a prescriptive set of obligations on actors involved in the development, deployment, and distribution of AI systems within the EU, mainly the providers, deployers and importers. These obligations are structured around a risk-based classification system that determines the level of compliance required: unacceptable risk, high risk, limited risk and minimal risk.

AI systems deemed to present an unacceptable risk to fundamental rights, safety, or the rule of law are simply prohibited. There are eight prohibited practices, including: biometric categorization that deduces or infers sensitive attributes (e.g., political opinions, religious beliefs, sexual orientation); subliminal, manipulative, or deceptive techniques that materially distort an individual's or group's behavior, leading them to make decisions they would not otherwise make and causing harm; and real-time remote biometric identification in publicly accessible spaces by law enforcement.

High-risk AI systems are subject to the most stringent requirements. This category comprises: (1) AI systems used as products or safety components of products that need to undergo conformity assessment under EU law (e.g., machinery, toys, radio equipment) and (2) standalone AI systems in specified use-cases (e.g., biometrics, critical infrastructure management, decision making in recruitment and human resources, law enforcement uses).

Providers of **high-risk AI systems** must meet stringent ex-ante and lifecycle obligations, including: (i) Implementing and maintaining a continuous risk management system throughout the system's lifecycle; (ii) ensuring that training, validation, and testing datasets are relevant, sufficiently representative, and, to the extent possible, free of errors and complete for the intended purpose; (iii) preparing and maintaining comprehensive technical documentation covering all aspects of the system; (iv) designing systems to automatically generate logs to facilitate traceability and oversight, including logs for substantial modifications; (v) providing clear instructions and information necessary for deployers to use the system in compliance with the EU AI Act; (vi) designing systems to allow deployers to implement human oversight; (vii) ensuring appropriate levels of accuracy, robustness, and cybersecurity; (viii) undergoing a conformity assessment procedure, issuing an EU declaration of conformity, and affixing CE marking where applicable; and (ix) reporting serious incidents to relevant authorities and conducting investigations, including risk assessment and corrective actions.

Limited-risk systems are subject to targeted transparency obligations. Depending on the system, these obligations can include informing individuals that they are interacting with an AI system, signaling the use of emotion-recognition or biometric categorization, and labeling AI-generated or manipulated

content subject to few exceptions. In the current version of the text, the wide definition of “manipulated content” allows for a very broad interpretation that can place basic content editing features within the remit of these transparency obligations. Such uncertainty will need to be addressed in the Guidelines and a Code of Practice on Transparent Generative AI Systems that the European Commission is in the process of drafting¹⁴.

Minimal risk systems, which do not fall under the categories above, such as AI-enabled video games or spam filters, are not subject to binding obligations under the Act.

Finally, the EU AI Act lays down a dedicated framework for general-purpose AI (“GPAI”) models. The providers of these models are required to comply with specific obligations, including preparing and maintaining technical documentation; establishing a policy to comply with the EU law on intellectual property and related rights; and publishing a detailed summary of the data used for model training. When the GPAI models present systemic risks, providers are subject to additional obligations, such as assessing and mitigating systemic risks, conducting model evaluations and reporting serious incidents.

1.1.2. Capacity to adapt to the evolving AI ecosystem

The question of whether the EU AI Act’s provisions can adapt to future AI developments has long been debated and remains ongoing. This legislative challenge, which is not unique to AI regulation, has been illustrated in the legislative process leading to the adoption of the EU AI Act. Notably, the provisions addressing GPAI were introduced at a very late stage, as the European Commission’s initial 2021 proposal did not foresee the widespread deployment of large language models.

Some requirements may, indeed, not be well suited to technological evolutions, such as the emergence of agentic AI. AI agents pursue complex goals, act proactively rather than merely producing outputs in response to prompts, learn and adapt in dynamic environments. They are therefore highly autonomous. In this context, the application of certain provisions could be challenging, in particular the requirement of appropriate accuracy. This requirement presupposes the existence of a defined reference standard for the task performed by the AI system and assumes that risks can be identified by measuring deviations of the system’s outputs from that standard; assumptions that are difficult to reconcile with the operation and nature of agentic AI¹⁵.

Beyond adaptability to future AI developments, critics highlight that the EU AI Act is insufficiently aligned even with the current state of technology and that certain requirements are impractical, if not impossible, to comply with in practice. For instance, the obligation to ensure that datasets are relevant, sufficiently representative, and, to the extent possible, free of errors and complete has been criticized as conceptually broad and difficult to comply with.

Relatedly, the EU AI Act’s allocation of obligations across defined roles in the AI value chain, such as providers, deployers, importers, and distributors, assumes a relatively stable supply chain. However, in a rapidly evolving technological environment, value chains can become increasingly complex and

¹⁴ European Commission, [Draft Code of Practice on Transparency of AI-Generated Content](#), December 2025

¹⁵ Z. Meyers, D. Schnurr, P. Larouche, [The AI Act and Technological Neutrality](#), November 2025

dynamic. As a result, role classification and the allocation of responsibilities may prove difficult to apply consistently and effectively in practice¹⁶.

That said, the EU AI Act is not entirely devoid of mechanisms to accommodate technological evolution. In particular, the use of delegated acts empowers the European Commission to adjust certain aspects of the regulatory framework in response to technological developments, thereby seeking to render the EU AI Act future-proof to the extent possible. From this perspective, some stakeholders argue that the primary challenge lies not in the legislative framework itself, but in regulators' capacity to enforce it in a rapidly changing technological landscape¹⁷. Yet, many observers question whether regulators currently possess the necessary technical expertise and resources¹⁸.

1.1.3. Legal uncertainties

Although it is intended to establish a proportionate framework, the EU AI Act's risk-based architecture triggers legal uncertainty.

The risk-based approach leans heavily on intended purpose and a narrow list of use-cases, leaving many real-world AI services only partially covered, or not covered at all. Without a clearer pre-categorisation by sector (for example: financial, legal, social, health), French and Indian providers struggle to determine whether their AI tools fall into "high-risk" "limited risk" or out of scope category (low risk). That uncertainty cascades into product design decisions, documentation planning, and time-to-market. In addition, the use of narrowly defined use cases through the Annex III tied to specific services and their target users create a fragmented regulatory framework that fails to capture the broader dynamics of the AI ecosystem. Such a case by case logic prevents the establishment of coherent rules and conflicts with the foundational legal principle of moving from general norms to specific applications. As a result, it limits the capacity of the law to provide useful future proof guidance and undermines the overall effectiveness of AI governance.

This uncertainty is further amplified by the Digital Omnibus on AI of 19 November 2025, which introduces a "stop-the-clock" mechanism for high risk AI obligations. Under the original EU AI Act, high risk system requirements were scheduled to apply from 2 August 2026, with later dates for systems embedded in regulated products. The Omnibus proposal disrupts this structure by making the start of those obligations conditional on the Commission confirming the availability of harmonised standards, common specifications or guidance — effectively pausing the countdown until "adequate measures in support of compliance" exist but providing no clear legal test for what adequacy entails. What appears, on its face, to be a pragmatic flexibility tool in fact generates a new layer of unpredictability: deadlines no longer flow from the regulation itself, but from future administrative determinations whose timing stakeholders cannot anticipate. Providers may gain time, but lose the regulatory certainty required for investment planning and product sequencing. The Omnibus introduces long stop dates (extending

¹⁶ Z. Meyers, D. Schnurr, P. Larouche, [The AI Act and Technological Neutrality](#), November 2025

¹⁷ Atte Ojanen, Johannes Anttila, Thilo H. K. Thelitz, Anna Bjork, [Governing rapid technological change: Policy Delphi on the future of European AI governance](#), December 2025

¹⁸ Atte Ojanen, Johannes Anttila, Thilo H. K. Thelitz, Anna Bjork, [Governing rapid technological change: Policy Delphi on the future of European AI governance](#), December 2025

obligations potentially to late 2027 or even August 2028) yet preserves the possibility that the Commission accelerates obligations earlier if it deems support instruments sufficient. This dynamic increases legal and operational risk for nonEU providers, particularly in countries such as France and India where companies rely on predictable EU timelines to coordinate global product releases. It also raises significant governance concerns: if the Omnibus is not adopted before 2 August 2026, high risk system obligations would automatically enter into force under the original AI Act, regardless of whether the necessary standards exist. This scenario risks imposing immediate, unachievable compliance duties, potentially triggering strong pushback from industry, national authorities, and even downstream deployers unable to meet obligations that depend on absent standards. In turn, the coexistence of two possible reality paths—one where the Omnibus delays obligations through the stop-the-clock mechanism, and one where the original deadlines apply—creates regulatory bifurcation. Businesses must plan simultaneously for contradictory timelines, which further erodes legal certainty and may cause premature compliance investment or strategic delay. By tying compliance to administrative milestones outside the control of providers, the Digital Omnibus on AI unintentionally exacerbates the very uncertainty it seeks to resolve.

1.1.4. Complex governance and enforcement mechanism

The EU AI Act establishes a multi-layered enforcement and governance model.

At the national level, the key actors include: designated notifying authorities responsible for designating and monitoring conformity assessment bodies; independent notified bodies which conduct pre-market conformity assessments for high-risk AI systems; market surveillance authorities tasked with supervising compliance and enforcement; and fundamental rights protection authorities empowered to act where AI systems cause violations to fundamental rights violations such as privacy and non-discrimination.

At the EU level, the AI Office within the European Commission is responsible for the supervision of the GPAI models across the EU. The AI Office is supported by the European AI Board, composed of representatives of Member States, which will have advisory functions to facilitate the consistent and effective application of the EU AI Act.

This structure creates a broad, and potentially overlapping, network of bodies and authorities, resulting in a fragmented and complex enforcement landscape under the EU AI Act¹⁹. Fragmentation occurs not only within the EU AI Act itself but also across national legal systems, as some Member States have appointed multiple market surveillance authorities. Stakeholders have raised concerns regarding potential divergent implementation practices across sectors and jurisdictions, especially in cases where AI systems operate in multiple EU countries or are based on large models subject to EU-level supervision²⁰. By comparison, the General Data Protection Regulation (GDPR), which relies on a less complex enforcement and governance framework, has nonetheless experienced significant fragmentation in practice, increasing compliance costs. Furthermore, in many countries, national authorities have not yet clarified how these new bodies will coordinate their roles with existing regulators such as data protection authorities, cybersecurity agencies, or sector specific supervisors. Because

¹⁹ European Parliament, [Interplay between the AI Act and the EU digital legislative framework](#)

²⁰ European Commission, [Staff Working Document Accompanying the Proposal for Digital Omnibus Regulation](#)

national roles and responsibilities are still being defined, organisations face uncertainty about who remains in charge of interpretation, how requirements will be enforced, and how oversight will be shared.

1.1.5. Regulatory design and textual density

The EU AI Act is notable for the density and granularity of its legal drafting, as it contains an extensive corpus of recitals, articles, annexes, paragraphs, and sub-points. Its textual density and granularity exceed other several key instruments in the EU digital framework: the EU AI Act contains 342 numbered paragraphs, compared with 270 in both the GDPR and the Digital Services Act (DSA), and a higher level of detail with 579 sub-points across paragraphs, definitions, and annexes, versus 316 in the GDPR and 250 in the DSA²¹.

This complexity is further amplified by the Act's regulatory architecture, which relies heavily on secondary legislation and complementary instruments to be adopted by the European Commission. These include delegated acts, implementing acts and guidelines, all of which expand the body of materials relevant to compliance. While these instruments are essential for clarifying, completing and updating the EU AI Act's provisions, their cumulative effect, combined with the already dense text of the EU AI Act, creates an intricate framework. This could pose significant operational challenges, particularly for small and medium-sized enterprises (SMEs) that lack dedicated compliance resources.

1.1.6. Integration within the broader legal framework

Throughout the text, the EU AI Act makes numerous references to other EU laws, such as the GDPR, the DSA, the Cyber Resilience Act, and the NIS 2 Directive, which highlights its deep integration into the existing regulatory framework. While these references aim to ensure consistency across EU legislation, they also increase complexity for stakeholders, who must interpret the EU AI Act in conjunction with these instruments.

Adding to this interpretive challenge, overlapping requirements often arise²². For instance, the EU AI Act introduces obligations on transparency and human oversight that mirror provisions under the GDPR: deployers of high-risk systems must notify affected individuals when such systems are used, and individuals have the right to request an explanation of the system's role in decisions with legal or similarly significant effects. Similarly, both the NIS2 Directive and the EU AI Act require the implementation of risk management systems, which apply cumulatively when essential or important entities under NIS2 provide or deploy high-risk AI systems under the EU AI Act.

In some cases, such overlapping requirements may create inconsistencies. For example, alert thresholds, such as serious AI incidents (AI Act), major ICT incidents (DORA), significant cyber incidents (NIS 2), data breaches (GDPR) and reporting pathways differ across the AI Act, DORA, NIS 2 and the GDPR, generating duplicate notifications and inconsistent escalation.

²¹ Theodoros Karathanasis, [The AI Act: Balancing Implementation Challenges and the EU's Simplification Agenda](#)

²² European Parliament, [Interplay between the AI Act and the EU digital legislative framework](#)

Viewed within the broader EU digital regulatory framework, actors across the AI value chain face a highly complex compliance environment, requiring adherence to extensive, overlapping or sometimes inconsistent requirements.

1.1.7. Impact on innovation and competitiveness

Compliance costs and administrative burden

The primary concern raised by stakeholders relates to the substantial compliance costs imposed by the EU AI Act. For high-risk AI systems and GPA models with systemic risks, these costs are particularly significant due to the extensive nature of the obligations and their continuous, ongoing character. Providers of such systems must navigate a complex regulatory framework, requiring considerable investment in both internal and external resources.

The impact assessment accompanying the proposal for the EU AI Act estimated that the theoretical maximum compliance costs and administrative burden is around EUR 10 000 for providers of high-risk AI systems that follow standard business procedures, plus EUR 3,000–7,500 to demonstrate that the requirements have been met²³. Stakeholders, however, anticipate materially higher costs. Although it is difficult to precisely estimate EU AI Act compliance costs at this stage, a comparison with the GDPR is instructive, as it has faced similar criticism despite establishing a less complex framework in terms of obligations and governance. Estimates indicate that GDPR compliance costs can reach up to EUR 500,000 for SMEs and up to EUR 10 million for large organizations²⁴. These compliance burdens have been associated with EU companies reducing data storage by approximately 26% and data processing by around 15% compared to US companies²⁵.

Delays in product deployment

The EU AI Act's ex-ante regulatory requirements introduce procedural steps that may materially slow the commercialization of AI technologies. Mandatory conformity assessments and documentation obligations prior to market entry can extend development timelines. These effects are amplified for AI systems and models that are frequently retrained or updated, because material changes can trigger re-assessment and documentation updates. In sectors characterized by rapid innovation cycles where time-to-market is a critical competitive factor, such delays may result in a loss of strategic advantage vis-à-vis jurisdictions with less restrictive regimes.

Chilling effect on innovation and competitiveness

Higher compliance costs and longer time-to-market, compounded with heightened liability exposure, are likely to stifle innovation and place European AI developers and researchers at a disadvantage vis-à-vis competitors operating under lighter-touch oversight.

According to the European Commission, the EU AI Act would classify only about 10% of AI systems as high risk. Yet several sectors are increasingly reliant on AI, such as healthcare where many innovations

²³ European Commission, [SWD\(2021\) 84 final – Impact Assessment Accompanying the Artificial Intelligence Act](#), April 2021

²⁴ Mario Draghi, [The Future of European Competitiveness](#), September 2024

²⁵ Demirer, M., Jiménez Hernández, D. J., Li, D., and Peng, S., [Data, Privacy Laws and Firm Production: evidence from the GDPR](#), February 2024

potentially fall within the definition of high-risk AI systems (e.g., predictive analytics and robotic surgery)²⁶. There is little debate that requirements such as risk management and human oversight are essential in a sector like healthcare. However, the associated burdens may disproportionately affect smaller players and advantage already dominant actors outside the EU, dampening innovation – particularly localized, needs-based innovation²⁷.

The EU's position in the global AI race is already precarious. Mario Draghi's report on the Future of European Competitiveness highlights that the key driver of the productivity gap between the EU and the US has been digital technology. Since 2017, 73% of foundational AI models have originated in the United States and 15% in China. In 2023, an estimated USD 8 billion in venture capital investment was made in AI in the EU, compared to USD 68 billion in the US and USD 15 billion in China. European companies developing generative AI models need substantial funding, needs that are not adequately met by EU capital markets, forcing many to seek financing abroad. Among the world's leading AI start-ups, only 6% of global funding flows to EU firms, highlighting the disadvantage the EU faces in attracting investment²⁸.

While this productivity gap is indeed due to multiple factors, regulatory approach plays a critical role. As Draghi notes, “the EU now faces an unavoidable trade-off between stronger ex ante regulatory safeguards for fundamental rights and product safety, and more light-touch rules to promote investment and innovation, for example through sandboxing, without lowering consumer standards.”²⁹ We note however, that regarding sandboxes³⁰, the EU AI Act creates a pro-innovation, rights-driven and supervised approach. As such Member States must establish at least one AI regulatory sandbox by 2 August 2026 (potentially in jointly with authorities of other Member States). Within these sandboxes, competent authorities must provide guidance, supervision, and support to identify and mitigate risks (especially to fundamental rights, health, and safety), and issue reports and written proof of completed activities that can help streamline later conformity assessment and market surveillance. These elements are designed to improve legal certainty and accelerate market access. Because the EU AI Act includes specific provisions in relation to sandboxes (in chapter VI), the EU model integrates sandbox participation with the EU AI Act's obligations, positioning the sandbox as both a compliance accelerator and a learning channel for market participants and regulators.

Other countries have taken different approaches to regulating AI and where the EU has opted for a horizontal regulation contained in a specific regulation, India has adopted a more flexible approach based on guiding principles and designed to encourage innovation. Drawing careful inspiration from this regime could help the EU sustain rights protection and safety without hampering innovation and sacrificing its competitiveness in the AI sector.

²⁶ Bignami EG, Russo M, Semeraro F, Bellini V, [Balancing Innovation and Control: The European Union AI Act in an Era of Global Uncertainty](#), October 2025

²⁷ Bignami EG, Russo M, Semeraro F, Bellini V, [Balancing Innovation and Control: The European Union AI Act in an Era of Global Uncertainty](#), October 2025

²⁸ Mario Draghi, [The Future of European Competitiveness](#), September 2024

²⁹ Mario Draghi, [The Future of European Competitiveness](#), September 2024

³⁰ Defined under the EU AI Act as: a “controlled framework set up by a competent authority which offers providers or prospective providers of AI systems the possibility to develop, train, validate and test, where appropriate in real-world conditions, an innovative AI system, pursuant to a sandbox plan for a limited time under regulatory supervision”. EU AI Act Art. 3.55.

1.2. AI Regulation in India

1.2.1. Approach

India's approach to AI regulation aims at encouraging innovation and adoption, while protecting individuals and society from the risk of harm caused by the development or use of AI. In this domain, India's primary objective is to leverage AI for economic growth and global competitiveness.

At this stage, India has adopted a policy-driven, principles-based and sector-led approach to AI governance, explicitly prioritizing innovation, adoption and agility over early-stage and technology-based prescriptive regulation.

Rather than introducing a standalone, horizontal AI statute, India's framework relies on policy instruments, voluntary measures, coordination mechanisms, and the empowerment of existing sectoral regulators (and regulations) to address AI-related risks within their respective mandates. The guidelines mentioned below highlight that existing laws, such as those regulating consumer and data protection but also criminal law, can be used to govern AI applications which is why no separate law is anticipated to be drafted (recognizing however that existing laws may need to be amended to encompass AI applications).

1.2.2. Principles

India's evolving AI regulatory roadmap was most recently set in November 2025 with the release of the government's note on India AI: Governance Guidelines (the "Guidelines")³¹. It builds seven guiding principles: trust in foundation, people first, innovation over restraint, fairness and equity, accountability, understandable by design, safety resilience and sustainability. In this respect, the Indian approach aligns with how the UK will regulate AI based on five principles and how Singapore will do so based on nine dimensions (further explained below). In addition, the principles themselves also align, highlighting the shared concern between jurisdictions for issues such as transparency/explainability, safety, accountability, and fairness.

Trust is the Foundation

Trust (the essential enabler of AI innovation, adoption, and effective risk mitigation) must be built into every stage of the AI value chain, spanning the underlying technologies, the organizations that develop and deploy AI, the supervisory institutions that oversee it, and the users who engage with AI tools.

People First

This principle embodies India's human-centric vision for AI governance. AI should be conceived and implemented to enhance individual agency and reflect the values of the communities it serves. In governance terms, this means ensuring that, to the fullest extent practicable, humans retain ultimate authority over AI systems, supported by meaningful human oversight that safeguards accountability.

³¹ Ministry of Electronics and Information Technology – Government of India, [India AI Governance Guidelines](#), November 2025

Innovation over Restraint

India's AI governance regime expressly places a premium on innovation and adoption rather than precautionary restraint. It treats AI as a foundational engine of economic growth, social advancement, and national development. Accordingly, governance should enable responsible experimentation and scaled deployment, ensuring that innovation is not stifled by premature or excessively restrictive regulation.

Fairness and Equity

In line with India's commitment to inclusive development, AI systems should be conceived, validated, and deployed to deliver outcomes that are fair, unbiased, and non-discriminatory, with particular attention to the needs and rights of marginalized and vulnerable individuals and communities.

Accountability

To preserve confidence in the AI ecosystem, developers and deployers must remain transparent and answerable for their roles. The guidelines emphasize that responsibility should be assigned explicitly according to each actor's function, the degree of risk presented, and the relevant due diligence requirements.

Understandable by Design

Understandability is a foundational design principle, not an afterthought. While AI systems are inherently probabilistic, they should, where technically feasible, offer clear explanations and disclosures that allow users and regulators to grasp how the systems function, their intended purposes, and the likely consequences of their deployment.

Safety, Resilience and Sustainability

AI systems should be designed with embedded safeguards that reduce the risk of harm and promote robustness and resilience. The guidelines emphasize the need for capabilities such as anomaly detection and early warning systems to anticipate, prevent, or mitigate adverse outcomes. As the official guidelines state: "The goal is to encourage innovation and adoption, while protecting individuals and society from the risk of harm caused by the development or use of AI." A sound governance framework reconciles these dual goals. Consistent with this approach, India generally focuses on regulating AI applications through sectoral regulators, rather than imposing rules on the underlying technology itself.

1.2.3. Recommendations

The Guidelines translate the seven guiding principles into a set of practical issues and policy recommendations, structured around six core pillars: infrastructure, capacity building, policy and regulation, risk mitigation, accountability, and institutional coordination.

Rather than proposing AI-specific binding obligations, the guidelines identify key systemic challenges (such as access to computing and data, skills gaps, fragmented oversight, emerging safety risks, and unclear responsibility allocation) and recommends incremental, proportionate responses.

These responses include targeted incentives and financing support, leveraging digital public infrastructure, strengthening human and institutional capacity (training programs for government

officials, civil servants etc. to understand AI technology developments), encourage the use of locally relevant datasets to support the creation of culturally representative models and applications, developing India-specific risk frameworks, clarifying accountability across the AI value chain based on function and risk, and relying on existing sectoral laws with targeted amendments where necessary.

Overall, the recommendations reinforce India's preference for a flexible, adaptive and innovation-oriented governance model, centered on guidance, coordination and capability-building rather than prescriptive ex ante regulation.

1.2.4. Regulatory implications

The guidelines do not call for a dedicated AI law at this stage, reflecting the view that existing legal frameworks (including data protection, consumer protection and sectoral regulation) can be leveraged and, where necessary, incrementally adapted through targeted amendments. Regulatory intervention is therefore expected to remain incremental, proportionate and context-specific, with a strong emphasis on guidance, voluntary commitments and regulatory capacity building, rather than ex-ante compliance-heavy obligations.

As such, the Guidelines indicate that the Government has started to review existing laws and regulations to identify potential gaps with a view to amend these laws and adapt them to AI applications. This includes for example a need to review (i) the Information Technology Act 2000 to clearly define the roles of various actors in the AI value chain (e.g., developer, user, deployer), (ii) the Digital Personal Data Protection Act to address questions such as whether the principles of collection and purpose limitation are compatible with how modern AI systems operate. In addition, committees (such as the Department for Promotion of Industry and Internal Trade's committee) have been tasked with addressing specific AI-related questions such as the copyrightability of works produced by generative AI systems or a "Text and Data Mining" exception to enable AI development.

For businesses and developers, this translates into a relatively permissive and innovation-oriented regulatory environment, characterized by flexibility and reliance on sectoral oversight and evolving policy expectations, rather than detailed AI-specific statutory obligations.

The Guidelines also emphasize the fact that voluntary measures would also closely align with India's pro-innovation approach as they would provide AI actors with the flexibility required to allow AI-related innovation to grow and adapt in the local social and cultural context. However, the Guidelines mention that the measures should be proportionate to the risk of harm which aligns with the EU risk-based approach. More specifically, the Guidelines provide that "low-risk applications may require only basic commitments such as transparency reporting and grievance mechanisms, whereas high-risk applications in sensitive sectors such as health or finance may require additional safeguards".

1.2.5. Sandboxes

In relation to sandboxes, the Guidelines made a recommendation to create regulatory sandboxes to enable the development of cutting-edge technologies provided these tests produce evidence with details of the risks observed and the safeguards applied.

Likewise, in the white paper on “Shaping the AI Sandbox Ecosystem for the Intelligent Age” (August 2025)³² the Ministry of Electronics and Information Technology stressed that “The need of the hour is to build secure and controlled environments that enable agile experimentation while maintaining trust, safety and public interest”. Therefore, the goal of this white paper is to propose a framework to guide the establishment and operationalization of AI sandboxes in India.

It is interesting to note that this white paper builds upon the definition and the regime of AI sandboxes under the EU AI Act to identify the key objectives of such environments (e.g., exit reports, facilities, personal data processing). It also highlights what enablers were put in place by sandboxes around the world including in the UK, the U.S. and Singapore (e.g., free access to Google Cloud’s high-performance GPU infrastructure for GenAI prototyping, offers tools for model validation, robustness testing and AI governance metrics).

Finally, the white paper concludes with a set of roles and responsibilities of the various stakeholders in establishing sandboxes in India. For instance, government and policy makers should “Facilitate a national registry of AI sandboxes to promote interoperability, documentation of learnings and cross-sector knowledge transfer”, start-ups and innovators should “Actively engage in sandbox pilots that align with their product–market fit and domain strengths”, and Industry and technology partners should “Contribute cloud credits, LLM APIs and compute resources through corporate social responsibility (CSR) or partnership models”.

1.2.6. Governance setup

While avoiding a centralized AI regulator, the Guidelines nevertheless envisage a strengthening of institutional coordination through the creation of dedicated bodies. In particular, they call for the establishment of:

- an AI Governance Group supported by a Technology & Policy Expert Committee, tasked with cross-government coordination, policy coherence, and oversight of responsible AI practices; and
- an AI Safety Institute, intended to support AI safety research, risk assessment, testing and evaluation, capacity building, and engagement with international AI safety initiatives.

This “whole-of-government” governance model seeks to address the cross-cutting nature of AI while preserving the role of sectoral regulators as the primary actors for supervision and enforcement.

Other countries have taken different approaches to regulating AI and where the EU has opted for a horizontal regulation contained in a specific regulation, the UK or Singapore have adopted a more flexible and vertical (sector-specific) approach designed to mitigate existing risks while preserving innovation capabilities and speed. Drawing careful inspiration from these regimes could help the EU sustain rights protection and safety without hampering innovation and sacrificing its competitiveness in the AI sector.

³² Ministry of Electronics and Information Technology – Government of India, [Shaping the AI Sandbox Ecosystem for the Intelligent Age](#), August 2025

1.3. AI Regulation in the UK

1.3.1. Approach

Recognizing that technology and innovation move at a very high speed, the UK proposes an agile, principle-based framework that regulates the use of AI by context and outcomes rather than creating prescriptive, technology-specific rules, with the goal of supporting innovation while addressing material risks and building public trust.

The government's approach is deliberately iterative and non-statutory at the outset and implementation will be through existing regulators (ICO, FCA, CMA, and Ofcom) with domain-specific expertise so as to avoid heavy-handed requirements that could stifle innovation and slow adoption.

1.3.2. Principles

The UK has identified five principles intended to be interpreted and applied on a non-statutory basis across all sectors by regulators within their area of expertise: (i) safety, security and robustness, (ii) appropriate transparency and explainability, (iii) fairness, (iv) accountability and governance, and (v) contestability and redress. These main five principles build on the Organisation for Economic Co-operation and Development (OECD) values-based AI principles.

As regulators apply these principles to address the risks specifically posed within their remit, the principles will add to existing regulation, and reduce friction for businesses operating within specific industries. In addition, regulators will be expected to clarify what compliance should be like in their domains, including by referencing technical standards, assurance techniques, and management systems to embed the principles into business processes.

The application of these principles will be proportionate and riskbased, allowing regulators to prioritize or de-emphasize principles depending on sectoral risk profiles and existing statutory frameworks. For example, some regulators may emphasize technical robustness and postmarket surveillance in safety critical settings, while others may foreground fairness, transparency, and redress in consumer facing markets.

Following an initial implementation period, the government anticipates, when parliamentary time allows and if monitoring shows legislation is needed to achieve effective implementation, introducing a statutory duty requiring regulators to have due regard to these principles.

The UK's framework is designed to achieve three objectives: (i) drive growth and prosperity by making responsible innovation easier and reducing regulatory uncertainty, (ii) increase public trust in AI by addressing risks and protecting our fundamental values, and (iii) strengthen the UK's position as a global leader in AI.

The framework rests on four pillars: (i) a functional, future-proof description of AI, (ii) a context-specific, use-based approach, (iii) cross-sectoral principles, and (iv) central support functions to ensure coherence, horizon scanning, and risk monitoring across sectors

1.3.3. Analysis

Definition of AI

With respect to the definition of Artificial intelligence, so far, the UK has taken a different approach than the EU as it has not provided an official definition of AI but rather defines it with reference to two characteristics of AI: adaptivity and autonomy. By designing its approach to address the challenges created by these characteristics, the UK aims to future-proof its framework against unanticipated new technologies that are autonomous and adaptive. On the other side, the EU has opted for an official, albeit wide definition of AI (“AI System” specifically): “a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments”.

Although the EU approach allows for certainty and predictability, there is a risk of over-inclusion of certain edge-case systems that maybe were not initially targeted by the EU AI Act. Under the UK’s more agile and contextualized framework, sector-specific regulators can adapt the principles to the specific risks (e.g., example of edge case within EU AIA and one that can be excluded).

However, the UK’s approach can adversely generate unpredictability for companies and fragmentation of enforcement by relevant regulators. In this respect, it should be noted that although the definition of AI is provided at the EU level, the enforcement of most of the EU AI Act’s provisions will be conducted by Member State authorities which still allows for a degree of variability across the EU.

Technology neutrality

To best achieve context specificity, the UK will empower UK regulators to apply the cross-cutting principles as they are best placed to conduct detailed risk analysis and enforcement activities within their areas of expertise. This approach could help innovation as regulators could assess appropriately the risk posed by AI in their remit.

Leverage existing regulators’ expertise

The UK’s approach to AI regulation is designed to be enforced through existing sectoral regulators, rather than a new AI regulator. In that way, domain expertise, proportionality, and speed can be preserved in supervisory and enforcement activities. As explained in the White Paper, an AI tool triaging customer emails will be overseen and enforced differently from an AI-enabled medical triage system, even if the underlying techniques overlap, because regulators will calibrate oversight to the sectoral context and outcomes at stake: “an AI-powered chatbot used to triage customer service requests for an online clothing retailer should not be regulated in the same way as a similar application used as part of a medical diagnostic process” (§45).

In contrast, the EU AI Act has created several AI-specific authorities tasked with the implementation of the EU AI Act (e.g., AI Office) but will also leave it to Member States to designate authorities in charge of the implementation (e.g., Market Surveillance Authorities (“MSAs”) in particular for high-risk AI systems). However, the EU AI Act mandates that Member States designate as MSAs authorities that are already in

charge of the implementation of specific regulations in certain areas (e.g., in vitro diagnostic medical devices, safety of toys).

Guidance

To ensure regulatory clarity and practicality, UK regulators are expected to publish guidance clarifying how the main principles align with existing laws within their remit, what specific compliance entails, and how those obligations will be monitored and enforced. Where a use case spans multiple regulatory frameworks, this should include joint guidance to minimize inconsistent or duplicative requirements for firms. For instance, the Financial Conduct Authority has issued an AI Update following the Government's response to the AI White Paper.

Building on the UK's White Paper commitments to AI regulation, the UK has also issued Initial Guidance for Regulators (February 2024) to set out the considerations that regulators may wish to have when developing tools and guidance to implement the UK's approach to AI regulation and in particular the five principles.

New Central Functions

To ensure the regulatory model works in practice, central government will put mechanisms in place to coordinate, monitor and adapt the framework as a whole without undermining the independence of regulators. Such functions include monitoring, assessment and feedback, as well as support coherent implementation of the principles, support innovators (sandboxes).

Pro-innovation approach

The UK approach is expressly designed to be pro-innovation by avoiding rigid, one size fits all rules and placing regulator discretion and proportionality at the center of implementation so that the requirements match context and risk.

More importantly, and this is an important difference with the EU AI Act, the non-statutory start allows rapid learning and adaptation as AI applications evolves, enabling the UK to respond to new opportunities and risks without imposing early burdens that could chill experimentation and adoption.

As the UK favors a regulator-led, sector-specific model, multiple regulators contribute to the establishment of sandboxes. For instance, in 2016, the UK launched the Financial Conduct Authority's fintech sandbox, and the Information Commissioner's Office has operated a regulatory sandbox to help organizations test innovative data-driven solutions in compliance with data protection requirements, often relevant to AI deployments.

More recently regarding sandboxes, it can be considered that the UK shifted from its initial regulator-based approach to a more horizontal approach.

On October 2025, the UK Government unveiled a new approach to AI regulation aimed at driving innovation, growth, and public trust. Central to this is the creation of AI Growth Labs acting as a cross-economy sandbox that will allow companies to test AI products in real-world conditions under controlled environments. As such, the Department for Science, Innovation & Technology indicated: "The Lab's cross-economy design will enable innovations which cut-across traditional regulatory boundaries to be tested".

Similarly (in part) to the sandbox regime under the EU AI Act, the AI Growth Lab would maintain public trust by monitoring sandbox participants under close and careful supervision.

1.4. AI Regulation in Singapore

1.4.1 Approach

Rather than a single binding AI law, Singapore provides detailed guidance and tools that organizations can operationalize, coupled with sectoral obligations (e.g., in finance) and baseline data protection requirements.

Singapore has developed a non-binding and principle-based approach to AI regulation, which differs from the risk-based approach taken by the EU AI Act. Minister Josephine Teo said that there are no immediate plans to introduce overarching laws to regulate AI. Singapore takes the view that existing laws already address harms associated with AI and, if necessary, they can be updated to address inadequacies.

For instance, regulatory guidance has been issued the Personal Data Protection Commission has issued a set of advisory guidelines (e.g., Advisory Guidelines On Use Of Personal Data In AI Recommendation And Decision Systems) on how the PDPA 2012 will apply at different stages of model development and deployment whenever personal data is used.

At the core of Singapore's approach is the Model AI Governance Framework, first published in 2019 by the Infocomm Media Development Authority ("IMDA") and subsequently updated to reflect the quick evolution of AI development. This framework aims at promoting the responsible use of AI by providing implementable measures for organisations to deploy AI.

1.4.2 Principles

Model AI Governance Framework (by IMDA)

The Model AI Governance Framework (2019, updated in 2020) (2020 Framework), provides detailed guidance to private sector organizations to address key ethical and governance issues when deploying AI solutions

In light of recent advances in generative AI, the AI Verify Foundation ("AIVF") and IMDA have also developed a draft Model AI Governance Framework for Generative AI (2024 Framework), which seeks to expand on the 2020 Framework by addressing new issues that have emerged from Generative AI and providing guidance on suggested practices for safety evaluation of Generative AI models.

This Framework promotes trust through 9 dimensions: accountability, data, trusted development and deployment, incident reporting, testing and assurance, security, content provenance, safety and alignment research and development, AI for public good.

AI Verify

AI Verify is a governance-focused testing framework and toolkit that enables organizations to assess their AI systems against 11 ethical principles using standardized evaluations. In practice, this tool allows

innovators to evaluate and generate a summary report to document the responsible AI practices implemented during the deployment of generative AI applications. This allows companies to be more transparent about their AI by sharing these testing reports with any stakeholders. In this respect, it can be noted that the EU has developed a comparable tool “EU AI Act Compliance Checker”. On this aspect India’s approach to AI regulation can align with Singapore’s approach as the Indian Guidelines have highlighted that integrating voluntary measures in the global AI regulatory framework could help the development of local AI-related innovation. However, India’s approach can be considered as going further as it suggests, to ensure that voluntary measures are widely adopted, implementing incentives (e.g., regulatory, financial). For instance, India’s Guidelines proposes to grant access to regulatory sandboxes to firms adopting voluntary safeguards, or to grant public recognition through certifications or ratings.

AI Verify was developed by Singapore’s Infocomm Media Development Authority (IMDA) in collaboration with private-sector stakeholders and is supported by the AI Verify Foundation (AIVF), a not-for-profit entity established by IMDA to harness industry and global open-source expertise in advancing AI testing frameworks, standards, and best practices. The IMDA has also submitted to public consultation a “Starter Kit for Safety Testing of LLM-Based Applications” as a set of voluntary guidelines collecting emerging best practices and methodologies for the safety testing of LLM-based applications. The Starter Kit focuses on four risks specific to LLM’s: hallucination, undesirable content, data disclosure and vulnerability to adversarial prompt.

Sandboxes

On 7 July 2025, the IMDA and the AI Verify Foundation introduced a Global AI Assurance Sandbox to create a testing ground for builders or deployers of Generative AI applications (excluding the underlying foundation models) to allow them to be tested by specialist technical testers. This initiative builds on the Starter Kit for Safety Testing of LLM-Based Applications as it emphasizes the need to consider the same four specific risks during testing. The sandbox carries several objectives including reducing testing-related barriers to adoption of generative AI, supporting the growth of a viable AI assurance market but also to informing policy guidance for the development of potential technical testing standards for GenAI applications.

1.5. AI Regulation in China and the United States

China

China’s AI regulatory regime has focused on particular use cases, such as generative AI and deep-synthesis technologies. According to one analysis, this approach reflects a vertical, technology-specific model that is influenced by national security imperatives alongside economic development goals.

The main regulation in China is the Interim Measures for the Management of Generative Artificial Intelligence Services (“Measures”), which entered into force on August 15, 2023. The Measures apply to entities that provide generative AI services to the public within China for the generation of text, images, audio, video, or other content, regardless of whether the service provider is based in China or abroad. They impose a set of substantive obligations, including requirements relating to privacy, content moderation, data labeling, fairness, and transparency. The measures follow a risk-based approach, although risk categorization is different than the one established under the EU AI Act. In particular,

generative AI service providers with "public opinion attributes or the capacity for social mobilization" are subject to heightened requirements, including security assessments and regulatory filing before launching the service.

Beyond the Measures, other important regulations include the Administrative Provisions on Algorithm Recommendation for Internet Information Services, which came into force on March 1, 2022 and the Provisions on Management of Deep Synthesis in Internet Information Service, effective January 10, 2023.

These regulations are supplemented by a growing body of guidelines and technical standards. Notable examples include the Practical Guidelines for Cybersecurity Standards – Method for Tagging Content in Generative Artificial Intelligence Services (2023) and the Basic Security Requirements for Generative Artificial Intelligence Service (2025), released by China's National Information Security Standardization Technical Committee.

United States

The approach to AI regulation in the United States reflects a combination of federal-level initiatives, state-level legislation, and voluntary industry standards rather than a comprehensive national regulatory framework.

At the federal level, President Joe Biden issued the Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence (Executive Order 14110) on 30 October 2023. EO 14110 directed more than 50 federal agencies and offices to undertake over 100 actions across 8 policy areas, including privacy and consumer protection. The Executive Order built on earlier federal efforts to promote responsible AI development, most notably the Blueprint for an AI Bill of Rights.

In January 2025, EO 14110 was rescinded by President Donald Trump through the Executive Order on Removing Barriers to American Leadership in Artificial Intelligence (EO 14179). Reflecting a deregulatory policy shift, this EO directs federal departments and agencies to review, revise, or repeal actions taken under the Biden administration "that are or may be inconsistent with, or present obstacles" to "America's global AI dominance".

In parallel with these federal developments, an evolving patchwork of state-level AI regulation has emerged, combining comprehensive risk-management laws with sector-specific and use-case-specific rules. The Colorado Artificial Intelligence Act, which will go into effect on February 1, 2026, adopts a risk-based approach to AI regulation that shares some conceptual similarities with the EU AI Act, focusing on high-risk AI systems. It imposes a series of obligations on developers and deployers of such systems, including requirements relating to documentation and transparency. California has pursued a more incremental approach, enacting multiple AI-related laws that address specific risks and applications rather than establishing a single comprehensive AI framework. These include rules governing automated decision-making and disclosure obligations for providers of generative AI systems. On December 11, 2025, President Trump signed an EO titled "Ensuring a National Policy Framework for Artificial Intelligence" which seeks to preempt state-level AI laws and regulations on the grounds that a fragmented regulatory landscape would impose significant compliance costs, while slowing US innovation in the global AI race.

2. Recommendations

In this international context, France and India share a vision that goes beyond regulation, and they intend to promote and advance this vision on the global stage. This common vision encompasses the principles that should guide the deployment of artificial intelligence within society, as evidenced by the joint declaration of 12 February 2025³³: safety, transparency, linguistic diversity and inclusive governance are established as values that should govern the implementation of AI systems and models.

This shared commitment to common values is all the more significant given that France and India have adopted different approaches to regulating artificial intelligence. Yet, these divergences should not be viewed as obstacles to meaningful cooperation and should not diminish the potential for collaboration, rather they underscore the complementary strengths each country brings to the partnership. Cooperation on AI need not rest on identical legal frameworks, but can instead be grounded in shared principles, mutual recognition of ethical imperatives, and a common commitment to harnessing AI for the benefit of society. By focusing on areas of convergence such as research and development, talent exchange, responsible innovation, and the promotion of AI applications in sectors of mutual interest, France and India can forge a partnership that transcends regulatory disparities and serves as a model for international collaboration.

Indeed, the foundations for such partnership have already been laid. France and India's cooperation on AI is rooted in the Horizon 2047 roadmap³⁴ published during Prime Minister Modi's visit to France in July 2023. This milestone was strengthened by the statement issued following President Emmanuel Macron's visit to India in January 2024. Both countries have continued to align their vision for AI governance, intending to safeguard fundamental rights while supporting innovation. This ambition has since taken concrete shape through various diplomatic engagements, including the co-presidency of the AI Action Summit in 2025, as well as President Emmanuel Macron's visit in India for the India AI Impact Summit in February 2026, in the context of which this White Paper has been prepared.

This trajectory of deepening engagement demonstrates that successful cooperation does not require regulatory uniformity. Looking back to the first steps taken around the world in terms of AI regulation reveals that the economic and ethical issues surrounding AI extend far beyond regulation alone. It is by embracing, but also moving beyond, sector-by-sector regulatory approaches, and the EU risk-based framework, that France and India can become major AI players on at the international stage.

To translate this shared ambition into tangible outcomes, concrete action is required. This section sets forth a series of practical recommendations designed to strengthen Franco-Indian cooperation on artificial intelligence.

Franco-Indian alliance as a driver of normative convergence

A Franco-Indian alliance on AI governance could serve as a cornerstone for global normative convergence. Rather than seeking strict legal harmonisation, the partnership would aim to develop a coherent set of foundational principles grounded in common values and shared geopolitical interests.

³³ [India-France Declaration on Artificial Intelligence](#), 12 February 2025

³⁴ [Horizon 2047: 25th Anniversary of the Indo-French Strategic Partnership: Towards a Century of French-Indian Relations](#), 14 July 2023

France and India, each with distinct legal traditions and regulatory approaches, are uniquely positioned to advance a governance paradigm that is adaptable, pluralistic, and inclusive.

The cooperation could focus on jointly promoting core ethical pillars such as:

- Non-discrimination and fairness
- Proportionality and responsible use
- Privacy and data governance
- Safety, robustness, and cybersecurity
- Transparency and explainability
- Human oversight and accountability
- Sustainable and responsible innovation

These dimensions reflect the principles highlighted in both the 2024 Joint Statement during President Emmanuel Macron's State Visit and the 2025 Franco-Indian Declaration on Artificial Intelligence. Embedding these shared commitments within a structured framework would signal a common vision for trustworthy AI and reinforce the credibility of both countries as global standard-setters.

Promoting Practical, Ethical, and Beneficial National AI Use Cases

For such principles to carry practical relevance, they should be supported by concrete, intelligible recommendations grounded in each country's domestic experience integrating AI for the public good.

India and France offer complementary and pioneering national use cases: India implemented AI in public administration early on, notably through the 2019 PM-Kisan programme, which automated benefit allocation to millions of farmers and demonstrated how AI can improve efficiency, reduce administrative errors, and enhance social equity. France, through initiatives such as Current AI, a \$400 million international programme dedicated to "AI for the public good", has emphasised the need to align AI innovation with societal benefit and international cooperation.

Drawing from these experiences, Franco-Indian guidelines could propose pragmatic orientations on:

- Inclusive research and design, ensuring marginalised communities and diverse linguistic groups are represented.
- Human-AI interaction, defining boundaries, safeguards, and user empowerment measures.
- AI for the public interest, enabling applications in health, agriculture, education, energy, and public services.
- AI and labour, promoting reskilling, decent work, and responsible automation.
- Safety, risk management, and resilience, applicable across high-impact systems.
- Innovation ecosystems, fostering responsible entrepreneurship and cross-border collaboration.

The overarching goal is to equip policymakers, developers, and civil-society actors with actionable tools to produce AI systems that are legally compliant, ethically aligned, and socially beneficial.

Establish a joint AI regulatory sandbox to enable collaborative testing of innovative AI

France and India should establish a joint AI regulatory sandbox to enable collaborative testing of innovative AI solutions under shared governance. This partnership would accelerate cross-border innovation and harmonize standards, reduce compliance friction while positioning both countries as leaders in responsible AI development.

Implement a coordinated, regulator-led approach for context-specific oversight

Implement a coordinated, regulator-led approach for context-specific oversight inspired by the UK model, leveraging existing sectoral regulators' expertise to calibrate requirements to use-case risk while maintaining cross-cutting principles of safety, transparency, fairness, accountability, and redress. In concrete terms, measures to clarify the risk-based approach by implementing a regulator-led approach could involve mapping the risk for each sector (financial, but also legal, economic, social, etc.) and abandoning the criterion of intended use of the AI system.

Maintain a technology-neutral, future-proof approach

Maintain a technology-neutral, future-proof approach to definitions and scope to prevent over-inclusion of edge cases while preserving legal certainty for stakeholders.

Foster partnerships between private-sector companies

France and India should foster partnerships between private-sector companies in key areas (e.g., health, finance, aerospace, banking, and automotive industries) to combine complementary strengths: France's advanced R&D and regulatory expertise with India's scale, cost efficiency, and growing tech ecosystem as well as both countries' advanced engineering skills. Such collaboration would accelerate innovation, open new markets, and position both countries as global leaders in next-generation AI-driven solutions.

Create a strategic alternative through a strong AI partnership between France and India

While the economic and technological dependence on the U.S. and China remains a recurring concern; establishing a strong AI partnership between France and India would create a strategic alternative, fostering innovation and reinforcing technological sovereignty for both nations.

France and India should take measures to enable secure and compliant data exchange between the two countries, as this would provide the foundation for joint AI initiatives, improve model training quality, and accelerate innovation across sectors.

Towards a global grammar for trustworthy AI

If structured effectively, a Franco-Indian governance initiative could evolve into a scalable model for international engagement. Beyond bilateral cooperation, France and India can contribute to the development of a global grammar for trustworthy AI – a shared conceptual and operational language that:

- Clarifies expectations for responsible AI design and deployment;
- Remains flexible enough to accommodate different legal systems and development levels;
- Encourages voluntary adoption rather than imposing formal treaty-level obligations;
- Supports norm diffusion through practical examples and demonstrable benefits.

Such a grammar would prioritise principles rooted in practice — such as transparency for high-impact systems, auditability and traceability of AI decision processes, and proportionate risk mitigation. Over time, these principles could strengthen mutual recognition, build cross-border trust, and support interoperable AI ecosystems, particularly in domains of public interest.

Ultimately, a shared commitment to human-centric and rights-preserving design can create the conditions for new and innovative technological players to emerge — offering high-performance AI services that are globally competitive, ethically grounded, and aligned with societal values.

While the AI Act constitutes a foundational step towards an “trustworthy and human centric” AI regulation, it remains limited in scope and cannot, on its own, address the transnational nature of AI development and deployment. As highlighted above, several regulatory gaps persist, particularly regarding cross-border cooperation, shared oversight practices, and the harmonization of ethical standards. As a result, the AI Act creates fragmented expectations for entities operating outside the EU territory.

As a consequence, Franco-Indian alliance in AI governance could also serve as a guide for global convergence through a shared articulation of general foundational principles. The main goal would be to settle normative coherence through values, overcoming formal legal harmonisation. France and India could therefore contribute to a governance paradigm that is compatible with diverse legal traditions. The governance could oversee non-discrimination, proportionality, privacy, safety, cybersecurity, innovation, transparency, and promote human oversight, which corresponds to the main pillars evoked both in the 2024’s Joint statement on the state visit of President Emmanuel Macron and the 2025 Franco-Indian declaration on artificial intelligence.

Conclusion

This first White Paper has set out a clear ambition: to turn the France-India relationship in artificial intelligence into a practical, structured and forward-looking cooperation agenda. Across the working groups, the contributions collected here converge on a shared diagnosis—AI is now a strategic capability that reshapes competitiveness, public services, and societal trust—and on a shared conviction: France and India have distinct, complementary strengths that, if combined with method and continuity, can generate outsized impact.

The recommendations formulated in these pages are deliberately action-oriented. They aim to reduce friction in collaboration, accelerate the translation from research to deployment, and strengthen the conditions of trust—through robust governance, validated solutions, responsible data use, skills development, and a sustained dialogue between policymakers, researchers, businesses and civil society. Beyond sectoral insights, the White Paper also highlights an essential message: in a fast-moving technological landscape, partnership is not a one-off statement of intent, but a long-term discipline—built on shared standards, tangible pilots, and mechanisms that enable scale.

The AI Impact Summit in India offers a timely and valuable anchor to maintain—and increase—this momentum. It provides a high-visibility platform to reconnect stakeholders, test early progress against real-world constraints, and ensure that France-India cooperation remains aligned with the pace of innovation and with evolving international discussions on safety, governance, and inclusion. In this context, the Summit is not a conclusion point, but a catalyst: a moment to consolidate what has been achieved and to open the next cycle of collective work.

Accordingly, the France India AI Initiative will continue to grow and to structure its actions over time. New occasions will be created to reconvene the existing working groups, deepen the recommendations, and translate them into additional proposals, pilots, and scalable pathways—drawing lessons from implementation and continuously raising the level of ambition. At the same time, new themes will progressively complement these preliminary works, reflecting emerging priorities, technological shifts, and the needs expressed by public and private actors in both countries.

Throughout this next phase, the Initiative will keep relying on what makes it distinctive: the instrumental involvement of the Young Leaders of the Foundation, and a broader France-India community that is consistently engaged, committed, and ready to build. Their energy, networks and capacity to bridge sectors will remain central to turning ideas into outcomes. The objective is straightforward: to ensure that France and India do not merely follow the trajectory of AI—but shape it together, through a partnership that is both strategic in vision and concrete in delivery.

CORPORATE PARTNERS

Thought & Content Partners France India AI Initiative



A&O SHEARMAN

BRUNSWICK

Partners Fondation France-Asie



LVMH



Galeries Lafayette

L'ORÉAL

TIKEHAU CAPITAL

KERING

SAINT-GOBAIN

SIMAERO



A&O SHEARMAN

BRUNSWICK



Partners France India Foundation



BNP PARIBAS



EXECUTIVE TEAM

Nicolas Macquin, Co-founder, President,
Fondation France-Asie,

Founder & CEO Archinvest

Arnaud Ventura, Co-founder,
Fondation France-Asie

Thomas Mulhaupt, Managing Director,
Fondation France-Asie

Agathe Gravière, Programs Associate,
Fondation France-Asie,

Hugues de Revel, Head of Programs,
Fondation France-Asie

Augustin de Buzonnière, Programs Associate,
Fondation France-Asie

Emma Lombard, Programs Associate, Fondation
France-Asie

Sumeet Anand, Co-founder, President,
France India Foundation,
Founder & CEO IndSight Growth Partners

Chiki Sarkar, Co-founder,
France India Foundation,
Founder, Juggernaut Books

Payal S. Kanwar, Representative Director General,
France India Foundation, Director General Indo-
French Chamber of Commerce & Industry

Anoushka Choudhary, Project Manager, France
India Foundation

Rishika Roy, Deputy Director, Northern Region,
Indo-French Chamber of Commerce & Industry

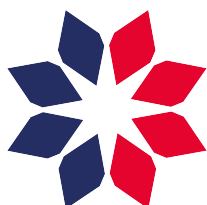
Yohann Samuel, Regional Director, Southern
Chapter, Indo-French Chamber of Commerce &
Industry

Shweta Pahuja, Regional Director, Western
Chapter, Indo-French Chamber of Commerce &
Industry

CONTACTS

thomas.mulhaupt@fondationfranceasie.org

payal@franceindiafoundation.org



**FONDATION
FRANCE-ASIE**

15 rue de la Bûcherie
75005, Paris France



**FRANCE INDIA
FOUNDATION**

B-38, Geetanjali Enclave
New Delhi, 110017, India





**FONDATION
FRANCE-ASIE**



**FRANCE INDIA
FOUNDATION**