



Social Security Scotland
Tèarainteachd Shòisealta Alba

Artificial Intelligence Adoption Framework

v. 1.1

August 2025

Dignity, fairness, respect.

Positioning

What is the Artificial Intelligence Adoption Framework?

The Artificial Intelligence Adoption Framework is a guide that advocates responsible, safe, and people-centric adoption of Artificial Intelligence in Social Security Scotland. It defines our principles, use cases, compliant data use, and risk appetite.

The framework is not a substitute for the existing lifecycle processes for new initiatives in the organisation. The framework complements existing processes and focuses on the management and governance of the opportunities and risks associated with Artificial Intelligence.

There is a need to formalise agreement and sponsorship of the AI Adoption Framework in response to the creep of Artificial Intelligence already present in Social Security Scotland.

Where does the Artificial Intelligence Adoption Framework fit?

Artificial Intelligence Adoption Framework

Initiate

- Proof of concept
- Initiative shaping
- Principle alignment
- Use case validation
- Data assessment
- Risk assessment
- Governance
- Architecture intent

Design and build

- High-level and low-level design
- Model procurement or build
- Data acquisition
- Re-use, procurement or development
- Testing
- Governance

Operate

- Service operations and continual improvement.
- Governance

The Artificial Intelligence Adoption Framework is founded on:

Social Security Scotland must adopt a position in respect of Artificial Intelligence (AI).

Where we adopt and leverage AI opportunities, we must maintain our people-centric focus.

The adoption horizon would transition from **assisted intelligence**, through **augmented intelligence**, towards a distant potential target state of **autonomous intelligence**.

Job roles will be paramount where implementation of Artificial Intelligence automates existing tasks and duties. Our people will be upskilled and trained to fulfil wider roles that deliver increased value, while enriching job satisfaction.

Artificial Intelligence is NOT Always the Correct Answer

The answer is Artificial Intelligence – but what is the question!

The use and adoption of Artificial Intelligence to resolve business problems and deliver client value must only be considered where appropriate, and alongside other available options.

Artificial Intelligence is not the “Silver Bullet” to every challenge.

Data Strategy Is a Mandatory Pre-Requisite

An active Data Strategy is a prerequisite to implementing the Artificial Intelligence Adoption Framework.

A Data Strategy that is sponsored, endorsed, and has active initiatives in place is a required foundation for implementing the Artificial Intelligence Adoption Framework.

The effectiveness of adopting Artificial Intelligence - in any capacity - will be dependent on the quality, integrity, availability and accessibility of Social Security Scotland's critical data elements.

AI Adoption Framework - Pillars

AI Principles

- Any initiatives that have an implicit dependency on Artificial Intelligence (AI) should adhere to our ethics informed principles to ensure these efforts are designed and implemented responsibly

AI Use Cases

- At the intersection of user needs and AI strengths
- Deliver demonstrable and measurable value to our clients

AI Data Sources

- Data is the fuel of AI
- Sourced and managed in compliance with our information and data policies

AI Risk Management

- Relevant to AI specific challenges in bias, discrimination, transparency and explainability

AI Governance

- Underpins and guides our obligations to policy, regulation and clients.

AI Principles

- **Inclusive Growth, Sustainable Development and Well-Being**
- Human Rights and Democratic Values, Including Fairness & Privacy
- Transparency & Explainability
- Robustness, Security & Safety
- Accountability

Guidance

"Stakeholders should proactively engage in responsible stewardship of trustworthy AI in pursuit of beneficial outcomes for people and the planet, such as augmenting human capabilities and enhancing creativity, advancing inclusion of underrepresented populations, reducing economic, social, gender and other inequalities, and protecting natural environments, thus invigorating inclusive growth, well-being, sustainable development and environmental sustainability."

- Inclusive Growth, Sustainable Development and Well-Being
- **Human Rights and Democratic Values, Including Fairness & Privacy**
- Transparency & Explainability
- Robustness, Security & Safety
- Accountability

Guidance

"AI actors should respect the rule of law, human rights, democratic and human-centred values throughout the AI system lifecycle. These include non-discrimination and equality, freedom, dignity, autonomy of individuals, privacy and data protection, diversity, fairness, social justice, and internationally recognised labour rights. This also includes addressing misinformation and disinformation amplified by AI, while respecting freedom of expression and other rights and freedoms protected by applicable international law.

To this end, AI actors should implement mechanisms and safeguards, such as capacity for human agency and oversight, including to address risks arising from uses outside of intended purpose, intentional misuse, or unintentional misuse in a manner appropriate to the context and consistent with the state of the art."

- Inclusive Growth, Sustainable Development Well Being
- Human Rights and Democratic Values, Including Fairness & Privacy
- **Transparency & Explainability**
- Robustness, Security & Safety
- Accountability

Guidance

"AI Actors should commit to transparency and responsible disclosure regarding AI systems. To this end, they should provide meaningful information, appropriate to the context, and consistent with the state of art:

- *to foster a general understanding of AI systems, including their capabilities and limitations,*
- *to make stakeholders aware of their interactions with AI systems, including in the workplace,*
- *where feasible and useful, to provide plain and easy-to-understand information on the sources of data/input, factors, processes and/or logic that led to the prediction, content, recommendation or decision, to enable those affected by an AI system to understand the output, and,*
- *to provide information that enable those adversely affected by an AI system to challenge its output."*

- Inclusive Growth, Sustainable Development Well Being
- Human Rights and Democratic Values, Including Fairness & Privacy
- Transparency & Explainability
- **Robustness, Security & Safety**
- Accountability

Guidance

"AI systems should be robust, secure and safe throughout their entire lifecycle so that, in conditions of normal use, foreseeable use or misuse, or other adverse conditions, they function appropriately and do not pose unreasonable safety and/or security risks.

Mechanisms should be in place, as appropriate, to ensure that if AI systems risk causing undue harm or exhibit undesired behaviour, they can be overridden, repaired, and/or decommissioned safely as needed.

Mechanisms should also, where technically feasible, be in place to bolster information integrity while ensuring respect for freedom of expression."

- Inclusive Growth, Sustainable Development Well Being
- Human Rights and Democratic Values, Including Fairness & Privacy
- Transparency & Explainability
- Robustness, Security & Safety
- **Accountability**

Guidance

"AI actors should be accountable for the proper functioning of AI systems and for the respect of the above principles, based on their roles, the context, and consistent with the state of the art.

To this end, AI actors should ensure traceability, including in relation to datasets, processes and decisions made during the AI system lifecycle, to enable analysis of the AI system's outputs and responses to inquiry, appropriate to the context and consistent with the state of the art.

AI actors, should, based on their roles, the context, and their ability to act, apply a systematic risk management approach to each phase of the AI system lifecycle on an ongoing basis and adopt responsible business conduct to address risks related to AI systems, including, as appropriate, via co-operation between different AI actors, suppliers of AI knowledge and AI resources, AI system users, and other stakeholders. Risks include those related to harmful bias, human rights including safety, security, and privacy, as well as labour and intellectual property rights."

AI Use Cases

- Assess **Assisted Intelligence** vs **Augmented Intelligence**
- Determine and evaluate the reward function
- Set expectations with all concerned – no silver bullets
- Determine the correct balance between control and automation (Human centric)

AI Adoption Framework – AI Use Cases – Positioning

The guidance on how to use AI is not intended to restrict or constrain innovative thinking. Instead, it is intended to ensure that careful consideration is given to the key principles and our overall plan.

Artificial Intelligence adoption in Social Security Scotland should be people-centric to the benefit of our clients and colleagues, with an understanding that fully autonomous AI is a long way off, and not on current horizons.

AI Adoption Framework – AI Use Cases – Human Centricity

Assisted Intelligence

Assisted intelligence helps colleagues perform tasks more effectively **without** learning from our data.

Expected to be our scope of use cases in the foreseeable horizons

Augmented Intelligence

Augmented intelligence provides further support to colleagues by **learning from our data**, understanding patterns, and enabling our people to rapidly complete complex tasks.

Autonomous Intelligence

Automating processes to deliver intelligence upon which machines, bots, and systems can act **independently** of human intervention.

Human Intervention

AI Adoption Framework – AI Use Cases – Drivers

It may seem obvious to deploy technology to solve business problems, but the use of Artificial Intelligence must show a clear connection to a **well-defined business problem**. The resulting solution should be measured to demonstrate its **value**.

AI Adoption Framework – AI Use Cases – Target Landscape

This is not a prescriptive direction, but rather guidance towards **potential** use cases that align with our principles and position. The following high-level flows do not represent the full scope of Social Security Scotland's AI Use Cases.

Client Engagement Flows

- Providing Information to clients
- Paying clients
- Receiving payment from clients
- Developing new policy
- Monitoring and evaluating existing policy
- Forecasting future demand

Enabling Business & Technology Workflows

- Procurement
- Legal and Compliance
- Technology Services

AI Adoption Framework – AI Use Cases – Target Landscape - Examples

Client Engagement Flows

- Providing Information to clients
 - Remove friction from information access
- Paying clients
 - Investigation and analysis
- Receiving payment from clients
 - Investigation and analysis
- Processing cases in the form of applications
- Expedite the benefit application process
- Enable benefit application to be transactional
- Ensure fair and equitable access to benefits
- Managing and sharing data
 - Remove friction from information access
- Developing new policy
 - Forecast demand for our services
- Monitoring and evaluating existing policy
 - Understand client opinion
 - Investigation and analysis
- Forecasting future demand
 - Forecast demand for our services

Enabling Business & Technology Workflows

- Procurement
 - Model and automate actions
- Legal and Compliance
 - Investigation and analysis
- Technology Services
 - Accelerate support tasks

AI Adoption Framework – AI Use Cases - A Focus on “Value”

- The wide-ranging potential of AI makes quantifying its value challenging; especially when combined with a perception that costs may be significant.
- There should be a focus on specific and measurable value that can be clearly articulated for Social Security Scotland, its workforce, and the people of Scotland – in alignment with Our Charter.
- Proactive and clear communication to emphasise value will be essential in building support and fostering trust.

AI Adoption Framework – AI Use Cases - Questions to Consider

ID	Question	Guidance
U01	What is the AI system or service being built or acquired and what type of product or service will it offer? Describe in terms of use case.	Expressed in terms of the business need and the explicit expected capabilities of the AI offering.
U02	Who are the users or clients of the proposed system or service?	If internal (specific division, branch & team impacted). If external (subset of clients to be defined if not all potential clients)
U03	What benefits will the system bring to its users and clients, and will these benefits be widely accessible?	This should be a value that can be measured. Identifying those that may be excluded from the value is as important as determining the target audience.
U04	Which part of Social Security Scotland or its suppliers—are responsible for building this AI system?	We should establish lineage of ownership. This is equally important internally or through vendor delivery. Vendor lock-in or reliance on a specific individual should be avoided.
U05	Which parts or elements of the AI system, if any, will be procured from third-party vendors?	It is key to understand the data that may have been used to deliver the outsourced component.

AI Adoption Framework – AI Use Cases - Questions to Consider

ID	Question	Guidance
U06	Which algorithms, techniques, and model types will be used in the AI system?	These should be explicitly declared with an assurance of their source being acknowledged.
U07	In a scenario where your project optimally scales, how many people will it impact in each target phase?	This should be represented as a series of target states; reflecting not only the additional impacted parties but the relative expected value and costs.
U08	Is the AI system's processing output to be used in a fully automated way or will there be some degree of human control, oversight, or input before use?	Fully autonomous would not honour our principles at this juncture. The design should enable a human in the loop or co-pilot approach.
U09	Will the AI system evolve or learn continuously in its use context, or will it be static?	For learning systems, we must establish clarity of how and where it will use our data to enable the learned outcomes.

AI Adoption Framework – AI Use Cases - Questions to Consider

ID	Question	Guidance
U10	To what degree will the use of the AI system be time-critical, or will users be able to evaluate outputs comfortably over time?	We should be facilitating the human in the loop maintaining control.
U11	What sort of out-of-scope uses could users attempt to apply the AI system, and what dangers may arise from this?	There is a need to look at the potential edge cases uses; including those that may be for detrimental purposes.
U12	Which business capability within Social Security Scotland will the AI system operate in?	There should be clear alignment to a business capability in keeping with the foundational need that the adoption must be related to a clear business demand.
U13	Have the subject matter experts within that business capability been consulted?	We need to show clarity of inclusion for the benefit of subject knowledge, and to mitigate fear that may accompany the adoption of AI.
U14	What is the value statement and business outcome?	Examples of these could be Drive efficiencies; improve quality; mitigate risks; reduce costs.

AI Data Sources

- Intent should be to gather high quality data from the outset
- Translate user needs (in context) to data needs
- Source the data in compliance with our policy – will restrict our use cases
- Social Security Scotland will assert full ownership, governance and control over bespoke models developed with our data

AI Adoption Framework – AI Data Sources – Positioning

It is widely acknowledged that the success of Artificial Intelligence is directly linked to the availability of **high quality, well governed** data.

For Social Security Scotland to adopt Artificial Intelligence we must conduct a thorough appraisal of our available data at the outset. The evaluation should consider accuracy, completeness, uniqueness, timeliness, validity, sufficiency, relevancy, representativeness, consistency and compliance.

AI Adoption Framework – AI Data Sources - Questions to Consider

ID	Question	Guidance
D01	What datasets are being used to build this Artificial Intelligence system?	There is a need to be explicit and to express clarity on content and ownership. No data sources without such lineage should be used.
D02	Will any data being used in the production of the Artificial Intelligence system be acquired from a vendor or supplier?	We need to understand where the vendor sourced this data and to establish its integrity and timeliness.
D03	Does the Artificial Intelligence system or service require that user data be transmitted off-site to a third-party system?	We must verify where data processing takes place. Social Security Scotland data should not be transmitted without ensuring data management processes are followed .
D04	Will the data being used in the production of the Artificial Intelligence system be collected for that purpose, or will it be re-purposed from existing datasets?	We should be particularly aware of our legal obligation in this context. Existing processes are in place to govern this – compliance should be demonstrated.
D05	Are there any anticipated data protection or intellectual property considerations about the data?	Existing processes are in place to govern this – compliance should be demonstrated.
D06	If the Artificial Intelligence system (or part of) is procured from a third-party, what is known about the data used to train the system?	We need to understand where the vendor sourced this data and to establish its integrity, bias etc

AI Risk Management

- Bias in AI based decisions
- Violating Personal Privacy
- Opacity and Misunderstanding in AI Decision Making
- Ambiguity in Legal Responsibility

AI Adoption Framework – AI Risk Management – Positioning

All existing risk management due process should be followed. We are not looking to replace the well-defined and understood protocols. However, it is imperative in the context of Artificial Intelligence, that we acknowledge the potential scale of harm that implementation may have.

Our approach is to advocate a viewpoint on the scale, scope and likelihood of the Artificial Intelligence adoption resulting in harm.

AI Adoption Framework – AI Risk Management – Context Approach

Scope

How many people could be adversely affected?

Scale

How severe could the harm be?

Likelihood

How likely is the harm to occur?

AI Adoption Framework – AI Risk Management – Context Approach

- Guidance

Context	Definition	Measures	Guidance
Scope	How many people could be adversely affected?	Calculated as a percentage of overall persons estimated to be affected by the Artificial Intelligence initiative.	The Artificial Intelligence Initiative team should establish a percentage threshold to determine a proportionate approach to demand stakeholder engagement and the definition of any required action. This will need to be balanced with scale as a high degree of harm to even just a few people may be an unacceptable degree of risk for us; given the service we are providing.

AI Adoption Framework – AI Risk Management – Context Approach

- Guidance

Context	Definition	Measures	Guidance
Scale	How severe could the harm be?	Catastrophic Harm: • Potential deprivation of the right to life; irreversible injury to physical, psychological, or moral integrity. • Deprivation of the welfare of entire groups or communities (those in receipt or entitled to benefits for example) • Catastrophic harm to democratic society or the rule of law • Deprivation of individual freedom and of the right to liberty and security. Critical Harm: • Significant and enduring degradation of human dignity, autonomy, physical, psychological, or moral integrity. • Enduring degradation of democratic society, or legal order.	We should assign a score from 1 to 4 where 4 is catastrophic. The score indicates the degree of proportionate response required. Any score greater than 1 is unlikely to accord with our Principles at this juncture.

AI Adoption Framework – AI Risk Management – Context Approach - Guidance

Context	Definition	Measures	Guidance
Scale	How severe could the harm be?	Serious Harm: - Degradation of human dignity, autonomy, physical, psychological, or moral integrity, or the integrity of communal life, democratic society, or just legal order or that harm to the information and communication environment Moderate or Minor Harm: - Does not lead to any significant, enduring, or temporary degradation of human dignity, autonomy, physical, psychological, or moral integrity, or the integrity of communal life, democratic society, or just legal order.	We should assign a score from 1 to 4 where 4 is catastrophic. The score indicates the degree of proportionate response required. Any score greater than 1 is unlikely to accord with our Principles at this juncture

AI Adoption Framework – AI Risk Management – Context Approach

- Guidance

Context	Definition	Measures	Guidance
Likelihood	How likely is the harm to occur?	Not Applicable: • It can be claimed with certainty that the risk of adverse impact does not apply to the adoption of AI in this initiative. Unlikely: • The risk of adverse impact is low, improbable, or highly improbable. Possible: • The risk of adverse impact is moderate; the harm is possible and may occur. Likely: • The risk of adverse impact is high; it is probable that the harm will occur. Very Likely: • The risk of adverse impact is very high; it is highly probable that the harm will occur.	We should assign a score from 0 to 4 where 4 is very likely. The score indicates the degree of proportionate response required. Any score greater than 1 is unlikely to accord with our Principles at this juncture

AI Adoption Framework – AI Risk Management

The decision to adopt any Artificial Intelligence solution should be evaluated within the context of its intended use. The evaluation should take into account our broader strategic direction and compliance specified within this framework. Proposed Artificial Intelligence solutions should be categorised using this structured classification:

Green

AI **adoption is appropriate**. Low risk of ethical or reputational harm. Anticipatory and ongoing human oversight and accountability protocols are in place.

Amber

AI adoption should be **treated with caution**. Potential risks have been identified. Defined mitigations and endorsement from empowered stakeholders is required for the adoption to proceed.

Red

AI adoption **should not proceed**. The identified and quantified adverse impacts outweigh any benefits.

AI Adoption Framework – AI Risk Management - Questions to Consider

ID	Question	Guidance
R01	Who are the stakeholders (including individuals and social groups) that may be impacted by, or may impact, the initiative?	We should have an approach that ensures there is no impact on those that have not been discretely identified.
R02	Do any of these stakeholders possess sensitive or protected characteristics that could increase their vulnerability to abuse or discrimination, or for reason of which they may require additional protection or assistance with respect to the impacts of the initiative? If so, what characteristics?	We must ensure compliance with Social Security Scotland's Charter in this respect as well as our wider legal obligations. Such an impact would likely deem this initiative as not appropriate to proceed.
R03	Could the outcomes of this initiative present significant concerns for groups of affected stakeholders with vulnerabilities caused or exacerbated by their distinct circumstances? If so, what vulnerability characteristics expose them to being negatively affected by the initiative outcomes?	We must ensure compliance with Social Security Scotland's Charter in this respect as well as our wider legal obligations. Such an impact would likely deem this initiative as not appropriate to proceed.
R04	What are the ethical considerations for this initiative in respect of sustainability?	We should be aware of the wider Scottish Government and UK Government objectives.
R05	What are the ethical considerations for this initiative in respect of safety?	Any initiative that poses a safety risk is likely to deem the initiative as not appropriate to proceed.

AI Adoption Framework – AI Risk Management - Questions to Consider

ID	Question	Guidance
R06	What are the ethical considerations for this initiative in respect of accountability?	We should be able to show transparency of process so that all stakeholders can understand how the initiative was conducted and why specific decisions were made.
R07	What are the ethical considerations for this initiative in respect of fairness?	Fundamentally this relates to demonstrating a clear understanding of the underlying data to acknowledge pre-existing bias, demonstrating a robust assurance approach and maintain “human in the loop”.
R08	What are the ethical considerations for this initiative in respect of explainability?	We should be aware of the balancing act here between the need for “showing the working” of the AI initiative; yet still protecting privacy and confidentiality.
R09	What are the ethical considerations for this initiative in respect of data stewardship?	Perhaps the most challenging. Data quality and Data integrity must be positioned to acknowledge the impacts on the Artificial Intelligence initiative.

AI Governance

- Internally; Executive Team, AI Working Group & Architecture Review Board
- Externally; AI Register

AI Adoption Framework – AI Governance – Context

Like other new initiatives, adoption of Artificial Intelligence initiatives should follow Social Security Scotland governance processes. However, AI presents unique challenges:

- **Evolution:** The pace of Artificial Intelligence advancements is expanding exponentially
- **Regulation:** Rules governing Artificial Intelligence are being developed which require compliance.
- **Ethical:** The use of Artificial Intelligence demands a strong focus on ethics. It is imperative that we ensure responsible use and mitigate potential bias. We must maintain our people-centric focus.

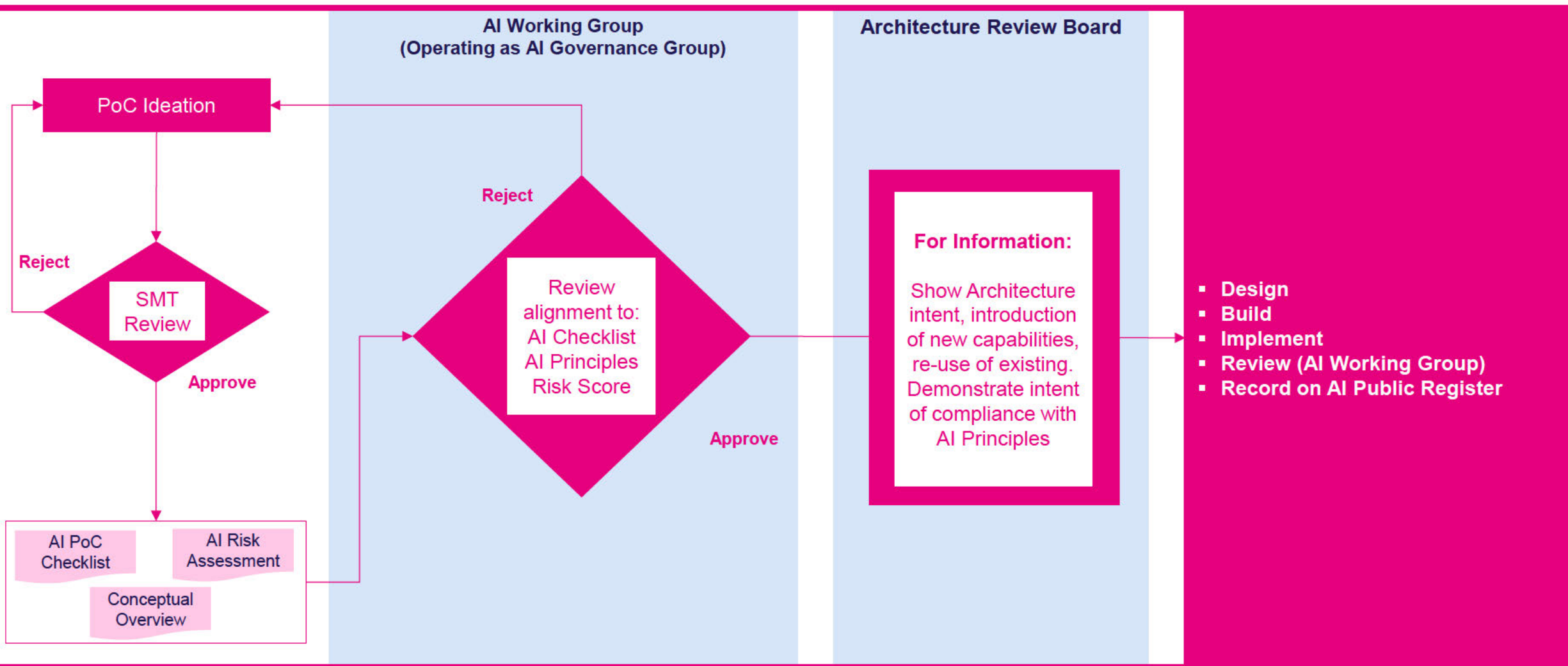
We do not recommend the creation of a separate “AI Project Board” for governance at present. Existing forums and formal entry gates to the Executive Team and the Architecture Review Board should be used.

AI Adoption Framework – AI Governance – Working Assumptions

- Activity streams that may generate the need for AI Adoption governance:
 - Proof of Concepts
 - Initiatives seeking Executive Team endorsement funding.
- Proof of Concepts will not proceed to operational stage without following the route to production governance.
- All Artificial Intelligence initiatives - Proof of Concept or Initiatives seeking ET endorsement funding - must be formally recorded on the AI Register.
- All AI related capabilities (including the ones already present in the products Social Security Scotland uses) should be governed through the new process ensuring alignment with the framework.

AI Adoption Framework – AI Governance – Proof of Concept

Proof of Concept (PoC)

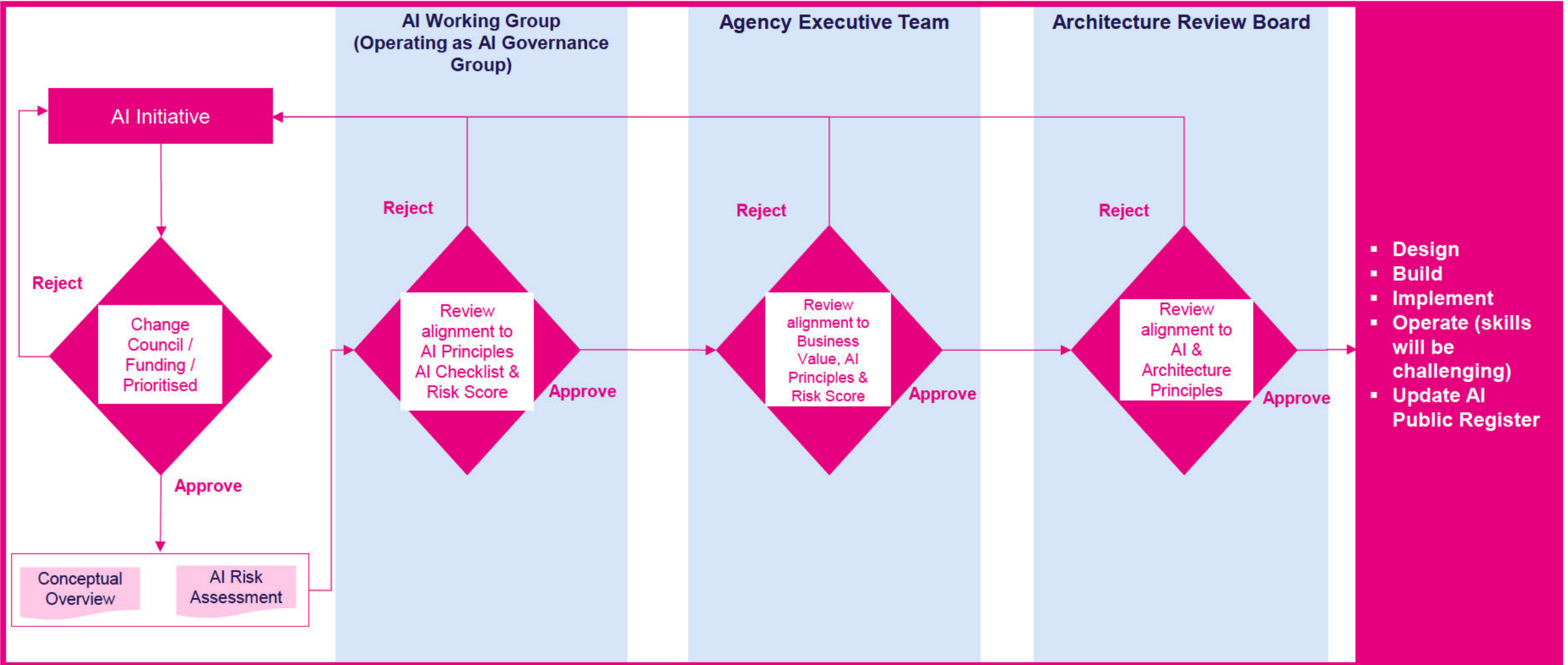


AI Adoption Framework – AI Governance – Proof of Concept

1. Any AI PoC must be checked to ensure that it has the characteristics required of any PoC (i.e. definition of what is to proven/disproven, statement of measures/evidence to be gathered to prove/disprove the concept).
2. SMT reviews the PoC proposal to ensure it is feasible and aligns with strategic direction.
3. The **AI Working Group**, acting as the AI Governance group, evaluates the PoC using the AI Checklist, the Principles in this framework, and Risk Score.
4. The **Architecture Review Board** is informed about the PoC and the introduction or re-use of existing capabilities.
5. Execution of approved PoC.
6. The **AI Working Group** reviews outcomes of the PoC and record it on the AI Public Register.

AI Adoption Framework – AI Governance – ET Endorsement

Initiative Seeking Executive Team Endorsement



AI Adoption Framework – AI Governance – ET Endorsement

1. Any AI initiative seeking endorsement must clearly define its purpose, expected outcomes, and alignment with strategic direction. Green light from the Change Council is required to proceed.
2. The AI Working Group (acting as the AI Governance Group) reviews the initiative for alignment with the AI Principles, AI Checklist, and assigns a Risk Score.
3. The Agency Executive Team assesses the initiative's alignment with Business Value, AI Principles, and Risk Score to determine strategic fit.
4. The Architecture Review Board evaluates the initiative for consistency with AI and Architecture Principles.
5. Upon approval, the initiative proceeds to design, build, and implementation stages. Note: operating skills may present challenges.
6. The AI Working Group updates the AI Public Register to ensure transparency and traceability.

Glossary

Term	Definition
AI Data Sources	The origin points of data used to train, validate, and operate AI systems, including structured, unstructured, internal, and external datasets.
AI Governance	A framework of policies, roles, responsibilities, and processes to ensure responsible, ethical, and compliant use of AI technologies.
AI Principles	Foundational ethical guidelines that guide the development and deployment of AI systems.
AI Use Cases	Specific applications or scenarios where AI technologies are implemented to solve business problems or enhance operations.
Artificial Intelligence	The simulation of human intelligence processes by machines, especially computer systems, to perform tasks such as learning, reasoning, and decision-making.
Assisted Intelligence *	AI systems that help humans perform tasks more efficiently, but do NOT learn from data or improve over time (e.g. GPS navigation).
Augmented Intelligence *	AI systems that collaborates with humans by learning from patterns and providing insights (e.g. financial portfolio advisors, AI for medical diagnosis).
Autonomous Intelligence *	AI systems that can make decisions and act on their own without human input (e.g. fully self-driving cars).
Bias	Systematic and unfair favouritism or prejudice in AI outputs, often caused by unbalanced training data or flawed algorithms.

Glossary

Term	Definition
Data Integrity	The accuracy, consistency, and reliability of data throughout its lifecycle, critical for trustworthy AI outcomes.
Data Strategy	A plan that defines how an organisation collects, manages, and uses data to support initiatives (including AI initiatives) and business goals.
Discrimination	Unjust or prejudicial treatment by AI systems toward individuals or groups, often because of biased data or algorithms.
Explainability	The degree to which an AI system's decisions and behaviours can be understood and interpreted by humans.
Model	A mathematical or computational representation trained on data to recognise patterns and make predictions or decisions.
People-centric	An approach to AI design and deployment that prioritises human needs, values, and well-being.
Reward function	A mechanism in machine learning, especially reinforcement learning, that defines the goal by assigning value to actions or outcomes.
Robustness	The ability of an AI system to maintain performance and reliability under varying conditions, including noisy or adversarial inputs.
Transparency	The openness and clarity with which AI systems, their data, and decision-making processes are communicated to users.