

Practicing Generative AI Responsibly

RESPONSIBLE AI DIMENSIONS · GENAI CHALLENGES · TRANSPARENCY TRADE-OFFS · SAFEGUARDS · SAGEMAKER & BEDROCK TOOLS · STUDY & REVIEW

01 MUST KNOW

if you remember five things, remember these

- 01 Responsible AI** = building and using AI systems that are **people-centric, ethical, safe, and unbiased** — with human oversight built into every stage of development and into production.
- 02 Eight AWS dimensions, three jobs:** **fairness, safety, privacy & security, veracity & robustness** protect people and data · **transparency + explainability** make AI open about how it decides · **controllability + governance** actively manage AI behavior in the organization.
- 03 Four generative-AI challenges:** **hallucinations** (plausible but false), **toxicity** (offensive output), **bias** (perpetuated or amplified from training data), **illegality** (IP or private data in outputs). Mitigable — never fully eliminable. No tolerance → use another approach.
- 04 Safeguards span the whole app:** prompt engineering + data curation → model transparency, explainability, evaluation → output filtering + human-in-the-loop → monitoring, auditing, feedback loops. **Data protection and governance stretch across everything.**
- 05 Tool map:** SageMaker **Model Cards** document · **Clarify** explains/evaluates (bias, toxicity, SHAP) · **Ground Truth** labels during development · **A2I** sends *production* predictions to human review · Bedrock **Guardrails** filters inputs and outputs.

02 KEY TERMS

TERM	DEFINITION
Responsible AI	Approach to building/using AI that is people-centric, ethical, safe, and unbiased
Fairness	No unequitable impacts on subpopulations (gender, ethnicity, age, language) — no unwanted biases
Veracity & robustness	Correct outputs even with unexpected or adversarial inputs and changing environments
Transparency	Degree to which you can understand how an AI model or system works
Interpretability	Understanding the <i>internal mechanics</i> — algorithms, data, decision process
Explainability	Clear, human-terms explanation of how a <i>particular output</i> was reached; model-agnostic
Hallucination	Output that sounds plausible but is verifiably untrue (e.g. fictional citations)
Toxicity	Output that is hateful, threatening, insulting, or demeaning to a person or group
Bias amplification	AI learning existing societal biases and perpetuating or magnifying them
Guardrails	Protective measures constraining AI behavior within safe, ethical boundaries (filters, checks)
Human-in-the-loop (HITL)	Integrating human oversight and decision-making into AI processes
SHAP	Shapley additive explanations — computes each feature's contribution to a prediction

03 THE EIGHT CORE DIMENSIONS OF RESPONSIBLE AI

AWS groups them by what they protect

DIMENSION	FOCUS	GROUP
Fairness	Minimize unequitable impacts on different subpopulations of users	Protect people & data
Safety	Avoid unintended harm to humans or the environment	Protect people & data
Privacy & security	Prevent data theft/exposure; individuals control use of their data	Protect people & data
Veracity & robustness	Correct outputs even with unexpected or adversarial inputs	Protect people & data

DIMENSION	FOCUS	GROUP
Transparency	Openness about capabilities, limitations, and risks	Be open about AI use
Explainability	Clear explanation of how a particular output was reached	Be open about AI use
Controllability	Human oversight + intervention keep AI aligned with human intent	Actively manage AI
Governance	Framework aligning AI with legal, ethical, societal standards	Actively manage AI

04 GENERATIVE AI CHALLENGES & SAFEGUARDS

know one example + one safeguard per risk

CHALLENGE	WHAT GOES WRONG	KEY SAFEGUARDS
Hallucinations	Plausible but verifiably false output — fictional citations, weird/creepy images; tolerance varies by use case	Fact-checking · deterministic inference parameters · data augmentation (RAG) · HITL validation
Toxicity	Discriminatory language, explicit material, harmful advice; boundary vs. censorship is context-dependent	Policies + guidelines · system guidance disallowing unsafe topics · exclusions and filters · human oversight
Bias	Societal biases in training data perpetuated or amplified — e.g. job descriptions assuming a male candidate	Bias-mitigation techniques · regular audits · continuous monitoring + feedback loops · human review
Illegality	Protected IP or private data in output — copyrighted passages, imitated celebrity voice, exposed PII	Automated plagiarism/IP checks · model auditing for leaks · HITL with SMEs and legal professionals

05 EXPLAINABILITY TECHNIQUES (MODEL-AGNOSTIC)

recognize each for the exam

TECHNIQUE	WHAT IT DOES
Feature importance analysis	Quantifies each input feature's contribution to the model's predictions
Partial dependence plots (PDPs)	Graph how changing one feature affects predictions, all other features held constant
SHAP	Calculates each feature's contribution to a specific prediction (Shapley values)
LIME	Builds a simplified, interpretable model around one prediction and explains that instead
Ask the LLM	Prompt the model to justify its own answer in natural language

06 SAFEGUARDS ACROSS THE GENERATIVE AI WORKFLOW

01 Prompt — engineering, data augmentation, curation → **02 FM** — transparency, explainability, evaluation → **03 Output** — content filtering, HITL reviews → **04 Ongoing** — monitoring, auditing, feedback loops → **05 Everywhere** — data protection + governance

07 AWS RESPONSIBLE-AI TOOLING

which feature does which job

TOOL	WHAT IT DOES	REMEMBER
SageMaker Model Cards	Standardized model documentation: details, performance metrics, intended use cases	Transparency + governance
SageMaker Clarify	Bias metrics + SHAP feature attributions; LLM evaluation for toxicity/accuracy — automatic or human jobs	The fairness/bias answer
SageMaker Model Monitor	With Clarify: detects bias drift + feature attribution drift in production	Drift, not demographics
SageMaker Ground Truth	Human labeling of text/image/video; results feed evaluation and fine-tuning	Development-time HITL
Amazon A2I	Routes low-confidence predictions to human reviewers; results in S3	Production HITL

TOOL	WHAT IT DOES	REMEMBER
Bedrock evaluation	Automatic or human-assisted jobs; toxicity, accuracy, robustness metrics	Also knowledge-base metrics
Bedrock Guardrails	Content filters · denied topics · word filters · sensitive-info filters · contextual grounding check	One guardrail, many FMs

08 EXAM TRAPS

where the knowledge check gets you

- ✗ “Interpretability and explainability are the same thing.” — **Interpretability** = understanding a model's internal mechanics (simple, transparent models). **Explainability** = a high-level, model-agnostic account of how an output was reached — the route for opaque-box FMs.
- ✗ “Maximum transparency is always the goal.” — Transparency **trades off** against accuracy, privacy, safety, and security — and can expose proprietary algorithms and model vulnerabilities to attackers.
- ✗ “SageMaker Ground Truth reviews production predictions.” — That's **Amazon A2I**. Ground Truth labels data and evaluates outputs during *development*.
- ✗ “Fix a pre-trained FM's bias by cleaning its training data.” — You **don't control** a pre-trained FM's training data. Data curation applies only when you customize or build; otherwise evaluate the model and mitigate bias at the outputs.
- ✗ “Content filters only screen model outputs.” — Filters apply to **inputs and outputs** — including catching prompt attacks and misuse attempts on the way in.
- ✗ “Model Monitor is the SageMaker feature for detecting demographic bias.” — **Clarify** detects bias and explains predictions; Model Monitor watches deployed models for drift and quality (Clarify integration adds bias-drift detection).
- ✗ “With enough safeguards, hallucinations are eliminated.” — Risk is **inherent** to generative AI. Safeguards reduce it; if the use case can't tolerate the residue, use a different approach.
- ✗ “The easiest control point is the training data.” — It's generally **less effort to filter the output** than to curate training data and prompts, given the breadth of tasks an FM supports.

09 SELF-QUIZ

answers are upside-down below

- Q1** Name the four principles that responsible AI encompasses.

Q2 Which two responsible-AI dimensions cover *actively managing* AI?

Q3 List the four responsible-AI challenges that generative AI amplifies.

Q4 Which type of bias arises from skewed human feedback during development or use?

Q5 Which explainability technique builds a simplified local model around a single prediction?
- Q6** Which Amazon Bedrock Guardrails feature helps verify the factual accuracy of output?

Q7 Which SageMaker feature provides standardized documentation of a model's details, metrics, and intended use?

Q8 Which AWS service sends a deployed model's low-confidence predictions to human reviewers?

Q9 Name two fairness metrics from the module (range 0–1).

Q10 Sample exam: which SageMaker feature is MOST effective for ensuring fair, unbiased predictions across demographic groups?

ANSWER KEY (rotate the page)

Q1 people-centric, ethical, safe, unbiased **Q2** controllability and governance **Q3** hallucinations, toxicity, bias, illegality **Q4** interaction bias **Q5** LIME (local interpretable model-agnostic explanations) **Q6** the contextual grounding check **Q7** SageMaker Model Cards **Q8** Amazon Augmented AI (A2I) **Q9** demographic parity, predictive parity, conditional use accuracy equality (any two) **Q10** SageMaker Clarify