



MOBILE TELECOMMUNICATIONS COMPANY (ZAIN)

RESPONSIBLE AI POLICY

Table of Contents

1 Purpose and application.....	3
2 Scope.....	3
3 The Responsible AI principles	4
3.1 Privacy and lawful data use	4
3.2 Security by design	4
3.3 Transparency and explainability	4
3.4 Fair treatment of customers/employees	5
3.5 Accountability and human oversight	5
3.6 Robustness, safety and accuracy	5
3.7 Continuous monitoring and improvement.....	6
3.8 Environmental sustainability	6
4 Prohibited uses	6
5 Risk-based governance	7
6 Data, privacy and residency	7
7 Generative, retrieval-augmented and agentic AI	8
8 Vendors and partners.....	8
9 Governance	8
10 Incidents and escalation.....	9
11 Training and awareness.....	9
12 Alignment with external frameworks.....	9
13 Review, ownership and coming into force	10
A Annex A — AI risk taxonomy	11
A.1 Customer harm and unfair treatment.....	11
A.2 Vulnerability and exclusion.....	11
A.3 Privacy, consent and data use	11
A.4 Telecommunications data sensitivity	11
A.5 Network, service continuity and safety	11
A.6 Security exposure specific to AI	11
A.7 Regulatory and contractual.....	11
A.8 Reputational.....	11
A.9 Model integrity and drift	11
A.10 Environmental impact.....	11
B Annex B — ISO/IEC 42001 cross-reference	12

1 Purpose and application

1.1 This Policy sets the Group rules for the responsible design, development, procurement, deployment, operation and retirement of artificial intelligence across Zain. Its purpose is to allow Zain to expand its use of AI with confidence. The controls in this Policy exist to protect customers, the network, Zain's data and Zain's legal position, and to do so in a way that is proportionate to risk, so that low-risk work proceeds quickly and scrutiny is concentrated where it is genuinely needed. In applying these controls, Zain acts consistently with its human-rights commitments and the UN Guiding Principles on Business and Human Rights and treats responsible AI as an enabler of digital inclusion across the markets it serves.

1.2 This Policy operates alongside Group policies on data protection, privacy, information security, procurement, enterprise risk and records. Where those policies already govern a matter, this Policy adds the AI-specific requirements only; it does not duplicate the underlying reviews. Operational detail and internal procedures supporting this Policy are set out in the Responsible AI Operating Standard maintained by the Group AI Center of Excellence and ratified by the Group Responsible AI Committee under Group governance.

1.3 Standing evidence expectation. For any AI system in scope, Zain must be able to show on request: who approved it, what data it uses, how it was tested, which controls apply, how it is monitored, and how incidents involving it are handled. The obligations in this Policy are written so that each of these six answers exists as a recorded artefact rather than as an assertion. For higher-risk and generative systems, these artefacts include the use case's impact assessment, model or system documentation such as a model card, and a post-market monitoring plan. For AI systems in service before this Policy's effective date, this expectation is met on the timeline set out at Section 13.3.

1.4 Status of this Policy. This Policy is an internal Group governance instrument. It states Zain's own requirements for the responsible use of AI and does not create contractual rights or obligations, or any right of action, enforceable by customers, vendors or other third parties.

2 Scope

2.1 Entities. This Policy applies to Zain Group and all Operating Companies and to Group functions. Vendors and partners are bound under Section 8 wherever they deliver AI to Zain or handle Zain data in doing so.

2.2 AI system. An AI system is any system using machine learning, generative AI, retrieval-augmented generation, or agentic capability that satisfies at least one of the following conditions: it is deployed or relied upon beyond its individual creator; it processes customer, personal, confidential or regulated data; it interacts with customers; it acts on the network; or it informs or executes commitments to external parties. This includes systems that are built, bought, co-developed or embedded within a wider platform, intelligent process automation with AI components, and network AI. AI that Zain operates to serve its own retail and enterprise customers in the course of Zain's business is within scope. The development or sale of AI products to external parties as a commercial offering is outside the scope of this Policy and is governed through the applicable commercial and product-governance arrangements. Zain inputs carry their obligations with them: where such an offering uses Zain data, Zain-developed models or AI assets, or Zain-controlled platforms, the use of those inputs remains

subject to this Policy's data, security and vendor requirements. Zain's obligations scale with the role it plays for a given system: where it builds, fine-tunes or substantially modifies a model it carries provider obligations, and where it procures and operates a model it carries deployer obligations.

2.3 Personal and experimental use. A tool that an individual builds or configures for their own productivity, using non-sensitive internal data, on which no other person or process relies, is not an AI system under this Policy. Such use is governed by the Zain AI/LLM Usage Policy and by approved-tools and acceptable-use rules. This exclusion ends, and the tool becomes an AI system requiring governance, at the moment any one of the following occurs: it is shared beyond its creator; another person or process relies on its output; it is connected to customer, personal, confidential or regulated data; or it is connected to a production system.

2.4 Lifecycle. Coverage runs across the full lifecycle: discovery, qualification, design, development, deployment, operation and retirement. AI is registered before build; assessed before design; and monitored in production.

3 The Responsible AI principles

This section states the eight principles that every AI system in scope must satisfy. Each is stated in two parts: what the principle means, and what it requires. The principles are derived from the main external reference frameworks — the NIST AI Risk Management Framework, ISO/IEC 42001, the EU AI Act's high-risk obligations and the GSMA Responsible AI Maturity Roadmap — adjusted for the realities of a telecom operator.

3.1 Privacy and lawful data use

Every use of data in an AI system rests on a documented lawful basis and a declared purpose. AI does not change Zain's data protection obligations; it raises the consequences of getting them wrong, because AI systems consume data at scale and can expose it in new ways.

This principle requires: a documented lawful basis and declared purpose for every data use; data minimization applied to features, prompts and logs; consent-dependent uses operating only where current consent exists; data quality and validation appropriate to the use case, so that AI decisions rest on accurate and trustworthy data; and data residency obligations identified before design and satisfied before deployment.

3.2 Security by design

AI systems are exposed to attack methods that conventional application security does not address, such as manipulation of the system through crafted inputs and extraction of sensitive information through model behavior. Security work for AI therefore covers these methods in addition to standard controls, at a depth proportionate to the risk of the use case.

This principle requires: AI-specific security review for high-risk use cases, covering the relevant threat categories; security testing scaled to risk and exposure; adversarial testing such as red-teaming and threat modelling for high-risk and generative systems and baseline security hygiene for every system, at every risk level, without exception.

3.3 Transparency and explainability

People are entitled to know when they are dealing with AI and to understand decisions that affect them. Explanation is proportionate to the audience and the consequence: a customer needs a clear business reason, a regulator needs a reviewable account, and an auditor needs the evidence trail.

This principle requires: disclosure wherever a system is customer-facing, with a clear route to a human; the ability to explain decisions that materially affect a customer in business terms appropriate to the audience; source citation for generative systems operating in sensitive domains; and a route by which a customer can seek human review of a decision that materially affects them

3.4 Fair treatment of customers/employees

This principle protects customers from unlawful and unjustified discrimination by AI systems. Where AI materially affects employees — for example in recruitment, performance management or workforce monitoring — the same protections against unlawful and unjustified discrimination apply. It prohibits discrimination on legally protected or sensitive grounds; the use of unjustified proxies for such grounds; the unjustified exclusion of vulnerable customers from essential access; and accuracy gaps across customer groups that cause harm.

This principle requires: testing of customer-affecting systems for these specific failure modes, with results recorded; screening of model features for unjustified proxies; and monitoring of outcomes at customer-group level where the risk warrants it; and design attention to accessibility for customers with disabilities.

For clarity, this principle does not restrict legitimate commercial differentiation. Zain differentiates offers, pricing and service levels on lawful commercial grounds, including customer value, and that is intended business practice. The principle governs how differentiation is done, not whether it is done.

3.5 Accountability and human oversight

Every AI system answers to a named person, and consequential decisions remain under effective human control. Oversight is not a formality: the accountable owner and operators must be able to review, override and stop the system, and no commitment binds Zain externally without human approval. For the purposes of this Policy, a decision is 'consequential' where it materially affects a customer, an employee, the network, or a binding external commitment; 'material' is calibrated in the Group AI Risk-Tiering Standard.

This principle requires: a named accountable business owner for every use case; human review, override and stop capability over consequential decisions; human approval before any binding external commitment takes effect; and oversight depth that scales with the system's autonomy.

3.6 Robustness, safety and accuracy

Systems perform reliably within their intended context and fail safely outside it. For a telecom operator the sharpest expression of this principle is automation that acts on the live network, where a failure is a service event, not a software bug.

This principle requires: defined performance baselines and fallback behaviour for every system before deployment; and, for automation acting on the live network, staged autonomy with simulation, blast-radius limits, tested rollback and human override.

3.7 Continuous monitoring and improvement

Responsible AI does not end at go-live. Models drift, data changes, and generative systems degrade as their sources age. The obligation is to detect this in production and act on it, and to feed what is learned back into the system and into Group posture.

This principle requires: production monitoring for drift, performance and fairness, and for grounding quality in generative systems; defined triggers for investigation and re-certification; and post-incident learning routed into the Group Register and into governance guidance where a pattern emerges.

3.8 Environmental sustainability

AI consumes energy, water and compute at material scale, and model choices drive that consumption. As Zain expands its use of AI, the environmental cost of that use is a design consideration weighed alongside performance and cost, not an afterthought.

This principle requires: consideration of energy, water and carbon impact in the design and procurement of AI systems; preference for the smallest model and least compute that meet the performance baseline; and attention to the environmental terms and efficiency of AI vendors and hosting arrangements.

4 Prohibited uses

This section defines the AI uses that Zain does not permit under any circumstances. The prohibitions are deliberately narrow and precisely defined, so that they do not capture ordinary business operations. Operational detail supporting these prohibitions and the process for handling permitted-but-restricted uses is set out in the Responsible AI Operating Standard.

- **Social scoring** of individuals — general-purpose scoring of people's behavior or characteristics outside a specific, lawful business relationship and purpose.
- **Mass surveillance**, defined as population-scale monitoring or tracking of individuals beyond what is required to deliver, operate and secure telecommunications services, and without a lawful basis. This covers, in particular, covert tracking and indiscriminate location surveillance outside license and lawful-intercept obligations. For clarity: maintaining consented, first-party customer profiles in a customer data platform, and behavioral segmentation conducted on a lawful basis for marketing and customer value management, are ordinary data processing governed by privacy law and are not surveillance under this Policy.
- **AI designed to manipulate, deceive or exploit vulnerability**, including manipulative interface design and covert influence on customers or employees.
- **Autonomous action on the live network or on customer accounts without the required safety controls** — simulation, blast-radius limits, tested rollback, human override and full audit logging.

- **AI making or communicating binding external commitments** — pricing, contractual terms, regulatory representations, service-level commitments or credit terms — that take effect without accountable human approval. For clarity, the automated presentation of offers, prices or service options within parameters approved by an accountable owner is not a binding external commitment under this prohibition.
- **Processing personal, regulated or public-sector data in an AI system in breach of lawful basis, purpose limitation or a residency mandate**, or permitting a provider to train on Zain data without explicit contractual consent.
- **Deploying or continuing to operate an AI system that has not been registered and risk-tiered** under Zain's Group AI governance.

Nothing in this Section restricts lawful-intercept, emergency-service, fraud-prevention or network-security obligations carried out under applicable law and approved internal controls.

5 Risk-based governance

5.1 Zain applies a risk-based governance approach to AI. Every use case is registered in the Group AI Register and assigned a risk tier before design begins, using the Group AI Risk-Tiering Standard. Three tiers determine the depth of governance a use case must clear before and after deployment. High-risk deployments require Group Responsible AI Committee approval at acceptance, at go-live, at each increase in autonomy, and at re-certification. Medium-risk deployments are approved by the Operating Company AI Office or the Centre under the delegated authority of this Policy. Low-risk deployments are registered and self-assessed by the accountable business owner and countersigned by the Operating Company AI Office. These tiers give effect to a defined Group AI risk appetite, approved under Group risk governance and reviewed as the business and technology change. For medium- and high-risk use cases a documented AI impact assessment is completed before deployment, covering the people and rights the system affects, its reasonably foreseeable misuse and downstream harms, and the mitigations applied; for high-risk systems it is refreshed on each material change.

5.2 Governance effort concentrates where risk concentrates. The Committee acts by exception: high-risk approvals, precedent-setting questions, exceptions and incidents. The Risk-Tiering Standard defines the scoring dimensions, weights, thresholds and mandatory escalation triggers, and is owned and calibrated by the Committee against portfolio evidence.

5.3 Restricted uses. A defined set of uses require Committee review and, where applicable, per-market legal confirmation before deployment, regardless of computed risk score. The Committee maintains the current restricted-uses list under the Responsible AI Operating Standard.

6 Data, privacy and residency

6.1 Every use of data in an AI system rests on a documented lawful basis, declared purpose and appropriate consent where required. Data residency obligations are identified before design and satisfied before deployment. Where a market mandates in-country processing of a data class, that mandate governs; where residency conditions vary, the strictest applicable

condition governs. Cross-border transfers are permitted only where a lawful transfer mechanism applies.

6.2 AI systems apply data minimization to features, prompts and logs. Data used for training, evaluation and monitoring is subject to the same protection standards as the source data. Zain does not permit providers to train on Zain data without explicit written consent. Data used to train or evaluate AI is documented for lineage and provenance, so that its origin and lawful basis can be traced.

7 Generative, retrieval-augmented and agentic AI

Generative AI systems, retrieval-augmented generation, and agentic AI carry additional risk profiles that require specific controls beyond the general principles. Zain applies minimum requirements to these classes of system covering grounding and source integrity, output evaluation before deployment in customer-affecting contexts, staged autonomy for systems that take actions, controls on intellectual-property and copyright exposure in generated outputs, protection of confidential and personal data against leakage through prompts and retrieval context, and defined human control points for agentic behavior. The operational specifications for these controls are set out in the Responsible AI Operating Standard and calibrated against evolving practice.

8 Vendors and partners

8.1 A significant share of Zain's AI is bought or embedded rather than built. Responsible AI obligations travel with the contract, and Zain retains the evidence and rights it needs from suppliers.

8.2 Vendor engagements involving AI include, as applicable: restrictions on data use and residency; an explicit prohibition on training on Zain data without written consent; documentation and transparency obligations covering the system's intended use, limitations and evaluation results; conformity evidence where the vendor builds to the EU AI Act, and accredited ISO/IEC 42001 certification where claimed; model-provenance evidence covering training-data legality, documented model origin and licensing terms, and screening for sanctions- and export-control-sensitive sourcing where relevant to Zain's markets; AI-specific security requirements; incident notification duties; audit and monitoring rights; and exit provisions including data return. These requirements flow down to sub-processors, model providers, plug-in and tool providers, and managed-service operators that materially contribute to the AI system or access Zain data.

8.3 Partner access to Zain systems and data is governed by Zain's existing information security and access management policies. Providers may not use Zain data to train models without written consent; Zain's models, prompts, evaluation sets and knowledge assets are treated as confidential information.

9 Governance

9.1 Three bodies carry AI governance at Zain. The Group Responsible AI Committee (GRAIC) is the Group authority on Responsible AI posture, approves high-risk use cases and

exceptions, classifies incidents, and ratifies the operational instruments issued under this Policy. The Group AI Center of Excellence operationalizes the Group's Responsible AI posture, prepares Committee submissions, enables Operating Company capability, and reports compliance. The Group AI Steering Committee holds strategic and portfolio authority for AI and decides the strategic and investment implications of Responsible AI constraints. For the highest-risk and precedent-setting matters, the Committee may draw on independent external expertise for advice and challenge. Named executive sponsorship for the Group's Responsible AI programme sits with the Group AI Steering Committee, which champions the programme and secures the resources and investment it requires.

9.2 Responsible AI decision authority sits at Group level. GRAIC is the sole body in the Group that sets or interprets Responsible AI posture. Operating Companies do not establish separate Responsible AI committees. Each Operating Company maintains an AI Office as the operational channel for enactment in its market, and its Head of AI Office attends Committee deliberations as standing observer.

9.3 GRAIC reports to the Group AI Steering Committee and, on risk matters, to the Group Risk Committee, and through them to the Board. The Board holds ultimate oversight of the Group's AI activity.

10 Incidents and escalation

10.1 AI incidents — failures, harms, near-misses, and risk crystallizations — are captured, classified and responded to under Committee-owned mechanics. Material incidents are classified by severity across regulatory, customer, financial, reputational, and network dimensions, with the highest severity governing the response.

10.2 Regulators are notified where legal or contractual obligations require, in accordance with applicable law. The Group Risk Committee receives severity-based reporting through the Group Chief Risk Officer as Co-Chair of GRAIC. Post-incident learning is routed into the Group Register and, where a pattern emerges, into ratified governance guidance.

11 Training and awareness

Zain trains its people to work responsibly with AI. Training is provided appropriate to role, from broad AI-awareness training for the workforce to specialized training for accountable business owners, developers, procurement, legal, risk and audit functions. Training addresses the principles at Section 3, the practical obligations at the working level, and the specific risks and responsibilities of generative AI use. Beyond training, Zain builds the specialist AI capability it needs and works to embed a responsible-AI culture, using change-management practices so that responsible-AI expectations are understood and acted on in day-to-day work.

12 Alignment with external frameworks

12.1 Industry alignment. Zain engages with the GSMA Responsible AI Maturity Roadmap as the industry-wide framework for responsible AI in telecommunications, and takes its maturity dimensions as reference points for the Group's Responsible AI programme. Zain's

substantive commitment to the Roadmap is progressed on a schedule aligned with the maturation of the Group's AI management system and the readiness of its evidence base.

12.2 Management system. This Policy positions Zain's AI activity on the ISO/IEC 42001:2023 AI management system framework. The Group's certification path is documented internally and gated to defined readiness milestones, with a target certification date set and reviewed by the Group Responsible AI Committee.

12.3 Intergovernmental baseline. Zain's principles and practice reflect the OECD AI Principles as the widely-adopted intergovernmental baseline for AI in society, and are consistent with the trustworthy-AI characteristics of the NIST AI Risk Management Framework and the high-risk obligations of the EU AI Act where they apply.

13 Review, ownership and coming into force

13.1 GRAIC owns this Policy. The Convenor of GRAIC, being the Head of the AI Center of Excellence, is its custodian. The Centre resolves routine questions of interpretation; contested interpretation is decided by the Committee.

13.2 This Policy is reviewed at least annually, and on material regulatory change in any operating market, a material AI incident, material change to the Group AI operating model, or at the request of GRAIC or the Group AI Steering Committee.

13.3 Coming into force. This Policy takes effect upon Board adoption. Existing and in-flight AI systems are brought into compliance under the standing lifecycle obligations of this Policy, on their next material change or on their next re-certification, whichever comes first. Detailed transitional provisions for the operating community are set out in the Responsible AI Operating Standard.

13.4 Prohibited uses at Section 4 apply in full from the effective date and are not subject to transition.

A Annex A — AI risk taxonomy

This annex defines the shared risk language used across the Group's Responsible AI framework. Every use case is screened against these categories at intake and assessment. The categories are stated with the failure modes most relevant to a telecom operator.

A.1 Customer harm and unfair treatment

AI affecting customer outcomes inaccurately, unfairly, manipulatively or without recourse — for example wrong billing guidance, unjustified offer exclusion or poor complaint routing.

A.2 Vulnerability and exclusion

Unjustified reduction in access, affordability or service quality for vulnerable, low-income, elderly, disabled, rural or small-enterprise customers.

A.3 Privacy, consent and data use

Use of data beyond lawful basis, consent, declared purpose or minimisation — for example call transcripts entering training data, or customer details entering prompts.

A.4 Telecommunications data sensitivity

Use of location, traffic, subscriber-identifier or network-signalling data in ways that engage licence obligations, lawful-intercept boundaries, or sector regulatory duties.

A.5 Network, service continuity and safety

Automation acting on the live network without simulation, blast-radius limits, tested rollback or human override — the operator failure mode where a model error becomes a service event.

A.6 Security exposure specific to AI

Attack methods that conventional application security does not address: prompt injection, model exfiltration, training-data poisoning, indirect access via retrieval sources.

A.7 Regulatory and contractual

Breach of applicable law or contractual obligation through AI deployment, including cross-border transfer restrictions, sector regulation, or supplier terms.

A.8 Reputational

Exposure to public, customer, regulator, employee or investor concern from AI deployment, including through peer-industry patterns even where no operational failure has occurred at Zain.

A.9 Model integrity and drift

Model behaviour departing from the intended envelope through drift, data change, source degradation in generative systems, or unintended emergent behaviour in agentic systems.

A.10 Environmental impact

Material energy, water, carbon or wider resource consumption from AI training, inference or hosting that is disregarded in design, procurement or model selection.

B Annex B — ISO/IEC 42001 cross-reference

This annex positions this Policy against the ISO/IEC 42001:2023 AI management system framework, indicating where each of the standard's core clauses is addressed within the Group's Responsible AI framework.

This Policy is designed to achieve conformance with ISO/IEC 42001:2023 when read together with the Group Responsible AI Committee Charter, the Group AI Risk-Tiering Standard, the Responsible AI Operating Standard, the Group AI Register and the assurance cycle, as those instruments are brought into effect.

Clause 4 Context of the organisation

The Group's context, the scope of its AI activity, and the interested parties whose expectations bear on it are addressed at Sections 1 and 2 of this Policy.

Clause 5 Leadership

Top-management commitment, AI policy direction, and the structure of accountability are established at Section 9 (Governance) and Section 13 (Ownership), and are given constitutional form in the Group Responsible AI Committee Charter.

Clause 6 Planning

Objectives, risks and opportunities are addressed at Section 5 (Risk-based governance) and given operational specification in the Group AI Risk-Tiering Standard.

Clause 7 Support

Resources, competence, awareness, communication and documented information are addressed at Section 11 (Training and awareness) and in the operational instruments maintained by the Group AI Center of Excellence.

Clause 8 Operation

Operational planning and control, AI system impact assessment and lifecycle management are addressed at Sections 5 through 7 of this Policy and specified operationally under the Responsible AI Operating Standard.

Clause 9 Performance evaluation

Monitoring, measurement, analysis, evaluation, internal audit and management review are addressed at Section 10 (Incidents), Section 13 (Review), and through the Committee's standing reporting cycle and Group Internal Audit interface.

Clause 10 Improvement

Nonconformity, corrective action and continual improvement are addressed at Section 3.7 (Continuous monitoring and improvement) and through the Committee's amendment and review mechanics.

The Group's ISO/IEC 42001 certification path is documented internally and its sequencing is aligned with the maturation of the operating instruments issued under this Policy.