

THE HAVEN

Operated by the Autistic Girls Network charity — 1196655 (England and Wales), SC054837 (Scotland)

The Haven: Responsible Use of AI Policy

Document	The Haven: Responsible Use of AI Policy
Version	04.26
Status	live
Last reviewed	29 April 2026
Next review	29 April 2027
Owner	Data Protection Officer
Approved by	Proprietor

Date: January 2026

Review Due: August 2026

Coordinator: DPO & DSL

Nominated Governor: Head & Safeguarding Committee

Version: 04.26

Consent-First Artificial Intelligence in Education and Care

Status: Organisational Policy

Applies to: All staff, contractors, volunteers, leaders, and technology partners

Scope: All AI-enabled systems used in educational, safeguarding, administrative, analytical, or support contexts involving children and young people

1. Purpose

This policy establishes a **rights-led, consent-first framework** for the use of Artificial Intelligence (AI) in education and care contexts.

It exists to:

- Protect the ****rights**, dignity, and developmental freedom of children and young people
> **
- Ensure organisational compliance with, and extension beyond, the ****EU AI Act**
> **
- Provide **clear operational guidance** for staff and leaders
- Prevent harm arising not only from misuse, but from **structural overreach** in relational technologies
- Set a **sector-leading Diamond Standard** for ethical, safe, and future-proof AI use

This policy recognises that **relational AI systems can cause harm even when technically compliant**, and therefore embeds **Consent Infrastructure** as a core design and governance requirement.

2. Legal and Normative Foundations

This policy is grounded in and extends the following frameworks:

2.1 Rights of the Child (Primary Authority)

All AI use must uphold the **UN Convention on the Rights of the Child (UNCRC)**, including but not limited to:

- **Article 3** – Best interests of the child
- **Article 12** – Right to be heard
- **Article 13** – Freedom of expression
- **Article 16** – Right to privacy
- **Articles 28–29** – Right to education that develops personality, talents, and abilities
- **Article 36** – Protection from all forms of exploitation

Where conflicts arise between efficiency, optimisation, or institutional convenience and children's rights, **children's rights prevail**.

2.2 EU AI Act (Compliance + Extension)

This policy meets and extends the EU AI Act by:

- Treating ****** all child-related AI systems as high-risk by default
> ******
- Requiring ****** human oversight, auditability, and contestability
> ******
- Prohibiting prohibited practices outright
- Adding **Consent Infrastructure safeguards** not explicitly required by the Act but necessary in child-centred contexts

Where national or international regulation is weaker than this policy, **this policy applies**.

3. Core Principle: Consent as Infrastructure

3.1 Beyond Procedural Consent

Consent is not treated as:

- a checkbox
- a one-time agreement
- silence or continued participation

Instead, consent is treated as **infrastructure**:

a set of structural conditions that must be preserved continuously within AI systems.

3.2 Properties of Infrastructural Consent

Any AI system used must structurally preserve:

- **Legibility** – it must be clear when interpretation, inference, or synthesis is occurring
- **Refusability** – refusal must be possible without explanation or penalty
- **Reversibility** – consent must be withdrawable with forward effect
- **Non-inferability** – silence, ambiguity, or presence must not be treated as agreement
- **Temporal agency** – users must control pace, duration, and depth of engagement

If a system cannot support these conditions by design, it **must not be used**.

3a. The Diamond AI Posture (Practice Layer)

The institutional **Diamond Standard** that follows in section 4 (Safety / Sovereignty / Symmetry / Stewardship) is operationalised by a relational **Diamond AI Posture** that staff carry into daily practice:

- **Work with AI.** AI is engaged as a collaborative tool — for accessibility support, drafting, summarising, pattern-spotting, removing friction for neurodivergent learners. Engagement is active, critical, and visible.
- **Do not offload decisions to AI.** Professional, safeguarding, pedagogical, behavioural, SEND, pastoral, and disciplinary decisions remain with the named human role-holder. AI may inform; it does not decide.
- **Do not defer entirely from AI.** Categorical avoidance is also a failure mode — it locks neurodivergent learners out of accessibility benefits and leaves staff under-skilled when AI is unavoidably in the environment. The right answer is engagement under boundaries.

The four institutional facets in section 4 (Safety, Sovereignty, Symmetry, Stewardship) are the structural conditions that make this practice safe. The three-action posture is what living within those conditions looks like, day to day.

This dual frame — institutional Standard + relational Posture — is the Haven's expression of consent-first AI in education and care.

4. The Diamond Standard

All AI systems must satisfy **all four facets** below.

4.1 SAFETY

AI must **reduce foreseeable harm** and must never replace human safeguarding judgement.

Requirements:

- No AI system may make autonomous safeguarding, disciplinary, or exclusionary decisions
- AI outputs must never be the sole basis for decisions affecting a child's rights
- All digital media (image, audio, video) must be treated as potentially synthetic
- Identity verification must never rely on a single signal (e.g. face, voice)

4.2 SOVEREIGNTY

Children and young people retain **authorship of their data, identity, and meaning**.

Requirements:

- AI must not extract trauma, infer diagnosis, or define identity
- Children must be informed, in age-appropriate ways, when AI is involved
- Opt-out routes must exist wherever legally and practically possible
- Children have the **right to be unread**: not constantly interpreted or profiled

4.3 SYMMETRY

Any system that interprets a child must itself be interpretable, contestable, and auditable.

Requirements:

- No black-box scoring, ranking, or behavioural prediction
- Clear documentation of data flows, inference logic, and limitations
- AI outputs require human corroboration
- **Forced coherence** (premature stabilisation of meaning) is treated as a safeguarding risk

4.4 STEWARDSHIP

AI exists to support human care and judgement — not replace it.

Requirements:

- Human accountability is non-transferable
- AI must strengthen relationships, not simulate them
- Systems must include **clear stopping conditions
> **
- Optimisation must never override dignity, care, or consent

5. Risk Classification (EU AI Act Aligned)

5.1 High-Risk by Default

Any AI system that:

- Processes children's data
- Interacts directly with children
- Influences assessment, placement, support, or behaviour
> is treated as **high-risk**, regardless of vendor classification.

5.2 Prohibited Practices (Absolute)

The following are **never permitted**:

- Biometric identification or categorisation of children
- Emotion recognition or affect inference
- Behaviour prediction or risk scoring
- AI-based webcam or biometric exam proctoring

- Training AI models on children's data
- Covert monitoring or surveillance

These practices are prohibited due to **structural coercion, unreliability, and rights violations**.

6. Allowed, Restricted, Prohibited Use

Allowed (Low Risk)

- AI supporting adults only
- Drafting, summarising, planning, admin
- No identifiable child data

Restricted (Formal Review Required)

- Any learner data
- Any child-facing AI
- Analytics, assessment, memory, or adaptation
- Requires AI Review + Safeguarding Approval

Prohibited

As listed above, regardless of vendor claims or efficiency gains.

7. Consent Infrastructure Safeguards

All approved systems must implement, where applicable:

- **Null zones** – spaces where no profiling or inference occurs
- **Non-fusion rules** – no implied shared identity or care
- **Refusal-preserving interfaces** – “no” without penalty
- **Temporal friction** – no instant conclusions
- **Containment & forgetting** – bounded memory and decay

8. Governance and Oversight

8.1 Roles

- **AI Governance Lead** – risk assessment, approvals, audits
- **Designated Safeguarding Lead (DSL)** – child-facing review and incident response
- **Senior Leadership** – risk appetite and final approval for high-risk systems

8.2 AI Review Process

All restricted uses require:

- AI Request & Review Form
- DPIA Digital Platforms (where applicable)

- Child Rights Impact consideration
- Ongoing review and re-authorisation

9. Vendor and Procurement Standards

Vendors must demonstrate:

- EU AI Act compliance
- Child-rights impact assessment
- No training on children's data
- Auditability and explainability
- Deletion, containment, and refusal mechanisms
- Incident reporting and redress pathways

Failure to meet these conditions results in **non-procurement**.

10. Incident Response

Any concern involving:

- misuse
- hallucination
- identity fraud
- deepfakes
- boundary violations

must be reported immediately to:

- DSL
- AI Governance Lead
- Data Protection Officer (where applicable)

11. Training and Competence

All staff must receive:

- Initial AI safety and consent training
- Annual refreshers
- Scenario-based safeguarding practice

12. Review and Accountability

This policy will be reviewed:

- Annually
- After any major incident

- When regulation or technology materially changes

Children's rights remain the **non-negotiable anchor**.

13. Closing Statement

Artificial intelligence is not neutral.

Relational systems shape meaning, behaviour, and possibility.

This policy exists to ensure that:

- **coherence never replaces care
> **
- **assistance never overrides agency
> **
- **technology never outruns consent
> **

The Diamond Standard is not about limiting innovation.

It is about **civilising it**.

14. Monitoring and Review

This policy will be reviewed annually by the DSL, DPO, and digital learning leads. Urgent updates will be made if risks emerge or if legislation changes.

15. Linked Policies

- Child Protection and Safeguarding Policy v01.26
- Data Protection, Confidentiality & Privacy Policy v10.25
- Behaviour and Regulation Policy v10.25
- Examinations Policy v10.25
- Prevent Duty Policy v01.26_
- Digital Consent Guidance for Families
- Children Missing from Education Policy v10.25

This PDF was generated from the published policy at policies.thehavenonline.school on 26 August 2026. The website is the authoritative, most current version of this policy.