

PUBLIC POLICY

Ethics, AI & Human Rights Policy

Principles and safeguards for human-centred AI in maritime operations, training and research.

Organisation: ELNAV.AI d.o.o.

Version: 2.0

Effective date: 30 July 2026

Policy owner: Founder & Chief Executive Officer

Review cycle: At least annually and after material legal, organisational or technical change

1. Purpose and scope

ELNAV.AI develops systems for bridge safety, active watchkeeping, speech-based verification, operational training, resilient navigation and maritime intelligence. This policy applies to all products, funded projects, internal research, pilots, demonstrations and partner integrations, including work performed by employees, contractors and project partners on ELNAV.AI's behalf.

The policy is informed by the Charter of Fundamental Rights of the European Union, the UN Guiding Principles on Business and Human Rights, Regulation (EU) 2024/1689 (EU AI Act), the GDPR and safety-management principles relevant to maritime operations. Legal classification and regulatory obligations are assessed for each intended use; this policy does not claim certification or regulatory approval.

2. Core commitments

1. **Human authority.** AI supports professional judgement and does not displace the authority or responsibility assigned to masters, officers, crew, trainers or shore management.
2. **Safety before convenience.** Release, pilot and integration decisions consider foreseeable misuse, degraded modes, false alerts, missed detections and over-reliance.
3. **Support, not surveillance.** Human-facing systems are designed around operational support, dignity, proportionality and just-culture principles.
4. **Privacy by design.** Local processing, restricted retention, non-identifying outputs and purpose limitation are preferred.
5. **Traceability.** We maintain documentation, logs and decision records proportionate to the development stage and risk.
6. **Honest claims.** We distinguish concepts, prototypes, pilots, operational tests and commercial products, and do not claim certification or safety outcomes without evidence.

3. Product and project boundaries

3.1 Aware Mate

Aware Mate is designed as active-watchkeeping support. It may generate local graded prompts from sustained patterns relevant to alertness, attention, activity and watch engagement. It is not a medical device, psychological assessment or employment-performance scoring system.

7. No identity recognition.
8. No emotion recognition.
9. No raw video ashore by default.
10. Advisory outputs with human review.
11. No automated disciplinary or employment decisions.

3.2 Speech-based bridge systems

Helm Order Monitor and related checklist or speech systems are intended to compare spoken operational exchanges with expected actions, procedures or sensor context. They are not designed to identify speakers or replace bridge-resource-management practice. Alerts and records must be interpreted in context by authorised personnel.

3.3 OSC Training

OSC Training supports scenario-based emergency leadership training, structured observation and debriefing. It does not perform medical or psychological diagnosis. The platform organises evidence and reporting without replacing trainer judgement, and participant records are not to be repurposed for unrelated employment decisions without a separate lawful and ethical basis.

3.4 RouteShield-PNT and navigation support

Navigation-support projects may cross-check position, time, route and sensor evidence and present advisory warnings. They do not replace approved bridge equipment, statutory watchkeeping duties or the navigator's responsibility. Development objectives and performance targets are not represented as achieved capabilities until validated.

3.5 Environmental and traffic intelligence

Environmental-sensing and maritime-traffic projects are assessed for ecological disturbance, data provenance, lawful use of vessel or location data, false-alarm consequences and the risk of unsupported enforcement conclusions. Outputs should support authorised decision-makers, not create automatic findings of wrongdoing.

4. Human oversight and operational responsibility

12. Each system has a defined intended user, purpose, operating context and limitation statement.
13. Human users must be able to understand the practical meaning of an alert or result and decide whether and how to act.

14. Where practical, systems provide acknowledgement, pause, override, fallback or safe-degraded behaviour appropriate to the use case.
15. Training and instructions address limitations, foreseeable misuse and the risk of automation bias.
16. ELNAV.AI does not present AI output as a substitute for compliance with COLREG, SOLAS, the ISM Code, company procedures or professional judgement.

5. Safety and risk management

Risk management is proportionate to the development stage and intended use. Concepts and prototypes are not held out as operationally validated systems. Before a pilot or release, the project owner records the intended purpose, users, data flows, known limitations, test evidence and acceptance criteria appropriate to that stage.

17. foreseeable hazards, misuse and human-factors effects;
18. false positives, false negatives, latency and loss of input;
19. degraded modes, cybersecurity and recovery;
20. alarm load, nuisance alerts and over-reliance;
21. change control, version traceability and regression testing;
22. incident, complaint and user-feedback handling.

6. Transparency and evidence

Technical documentation, user information and public claims must reflect the actual product stage and available evidence. Model or rule behaviour is explained at the level needed for safe use, without exposing security-sensitive or protected intellectual property.

Test results are reported with context, including configuration, duration, sample, operating conditions, limitations and the distinction between supervised testing and deployment. ELNAV.AI does not infer fleet-wide casualty reduction, insurance savings or regulatory compliance from a limited pilot.

7. Fairness, accessibility and bias

Human-facing systems are reviewed for performance differences that may arise from language, accent, lighting, skin tone, eyewear, facial characteristics, disability, sex, gender, age or other relevant factors. Testing and mitigation are performed when lawful, technically relevant and supported by an adequate sample.

Protected characteristics are not used to rank, discipline or profile individuals. Where a demographic factor is not relevant or cannot be assessed reliably, the limitation is documented rather than hidden.

8. Privacy, dignity and worker rights

Data collection must be necessary and proportionate to the stated safety, training or research purpose. The company avoids continuous identifiable recording where derived or local signals can meet the purpose.

ELNAV.AI does not design workplace systems to infer emotions or mental states for employment, disciplinary or performance-management decisions. Watchkeeping and training outputs must not be presented as proof of intent, character, competence or medical condition.

Deployments involving workers require clear information, defined access, retention and escalation rules, and meaningful consultation where law, contract or good practice requires it.

9. Prohibited or unacceptable uses

ELNAV.AI will not knowingly design, sell or configure its systems for:

23. social scoring or covert behavioural ranking;
24. identity recognition or biometric categorisation unrelated to a necessary, lawful and approved purpose;
25. emotion recognition for employment, disciplinary or performance decisions;
26. solely automated decisions that materially affect a person's employment or rights without lawful safeguards;
27. deceptive interfaces intended to manipulate users or conceal surveillance;
28. automatic enforcement conclusions based solely on an AI output;
29. claims of certification, approval, compliance or proven casualty reduction without evidence.

10. Governance and accountability

The Founder & CEO is accountable for this policy. Each project has an identified owner responsible for stage-appropriate ethics, safety, privacy and claims review. Significant issues are escalated to management and, where relevant, customers, partners, legal advisers, funders or competent authorities.

Suppliers and partners are expected to provide information needed to assess data provenance, model limitations, cybersecurity, licensing and human-rights risks. Contractual roles do not remove ELNAV.AI's responsibility for its own design and claims.

11. Concerns, review and contact

Employees, contractors, users and partners may raise concerns without retaliation through their project contact or info@elnav.ai. The policy is reviewed at least annually and after material changes in law, intended use, product architecture or evidence.

Document history

Version	Change
Version 1.0	Original public document.
Version 2.0	30 July 2026. Removed obsolete emotion-recognition and stress-alert claims; corrected the previous sustainability text placed under safety; added current product boundaries, stage-based risk management, support-not-surveillance rules, unacceptable uses and evidence discipline.

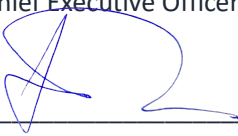
Approval

This policy is approved by the Founder & Chief Executive Officer and applies from the effective date shown above.

Approved by: Hrvoje Mihovilović

Role: Founder & Chief Executive Officer

Date: 30 July 2026

Signature:  _____