

UNI-PCSC Responsible Artificial Intelligence Policy (V2.45)

Approval Authority: Board of Directors / Sustainability Committee

1. Commitment & Vision

President Chain Store Corporation (UNI-PCSC) is dedicated to enhancing retail service experiences through digital innovation. We pledge that the development, procurement, and application of Artificial Intelligence (AI) will uphold our core values of Sincerity, Innovation, and Sharing, ensuring that AI systems are designed to enhance human well-being, promote social fairness, and foster environmental sustainability.

2. Governance & Accountability

2.1 Board Oversight: The "Sustainability Committee" and "Chief AI Digital Transformation Officer" are responsible for high-level oversight, designating personnel to be accountable for AI outcomes, risk decisions, and compliance. Board members commit to periodic training to ensure sufficient expertise in supervising AI strategies.

2.2 Responsibility Attribution: Every AI system must have a clearly designated System Owner to ensure legal accountability, this owner is responsible for the design, development, and operation of AI systems to protect personal data throughout the AI lifecycle. This includes compliance with data protection regulations, maintaining clear accountability paths for system deviations or harms, and assuming legal responsibility for AI-influenced decisions.

3. AI Authorized Use & Boundaries

To prevent AI systems from operating outside intended scopes and to ensure security, the company implements the following:

3.1 Authorized Scope: Clearly define AI functions and limitations. Usage is strictly limited to authorized business domains; any functional expansion leading to "Mission Creep" or unintended applications is strictly prohibited.

3.2 Allowed Use Policies: Establish normative guidelines to explicitly define what AI "can and cannot do" through both policy and technical means, translating these into specific technical restrictions.

3.3 Technical Guardrails: Implement system-level defense layers (Guardrails) and capability restrictions to mandate the filtering of inappropriate content and prevent AI from operating beyond its intended scope.

3.4 System Network Security: Adopt robust security measures to prevent cyberattacks, including penetration testing, incident response, encryption, and secure coding practices. This may also include mechanisms to prevent the unauthorized alteration of content.

3.5 Escalation Control Points: Human-in-the-loop (HITL) ensures that critical decisions influenced by AI include effective, documented human oversight, with authority and ability to intervene, override, or escalate. Oversight can range from human-in-command (final authority), to human-on-the-loop (monitoring), to human-in-the-loop (approval required). The system must ensure that AI systems cannot make irreversible or high-stakes decisions without human oversight and mechanisms for human review or override.

4. Environmental Sustainability & Resource Performance

Recognizing AI's environmental impact, the company takes the following Mandatory Measures to reduce its ecological footprint:

4.1 Supply Chain Mandate: We Mandate that own and third-party AI data center providers measure and report Energy use, Water use, and Carbon intensity. The use of energy-efficient infrastructure is mandatory, and the use of renewable energy is encouraged.

4.2 Resource Optimization: Mandate the implementation of energy-efficient infrastructure guidelines to reduce computational load and energy waste through algorithmic optimization.

5. Prohibited AI Practices

Strictly adhering to personal data protection laws and international privacy standards, the company prohibits the development or use of AI systems with the following human rights risks:

Manipulative Behavior: Prohibition of subliminal techniques or deceptive methods to distort human behavior.

Vulnerability Exploitation: Prohibition of exploiting specific groups (e.g., children, persons with disabilities) for inducement or exploitation.

Social Scoring: Prohibition of evaluating or classifying the trustworthiness of individuals based on social behavior.

Biometric Identification: Prohibition of unauthorized real-time remote biometric identification in public spaces.

6. Transparency, Fairness & Redress

6.1 AI Labeling: Proactively disclose the facts, purposes, capabilities, limitations, data sources, and risks of AI interaction. Provide understandable rationales for outputs and clearly inform users when they are interacting with AI rather than a human.

6.2 Bias & Drift Monitoring: Periodically conduct Bias Audits and detect "Model Drift". Implement AI designs that ensure fair and non-discriminatory behavior, including identifying biases in training data and ensuring the creation of diverse, representative datasets.

6.3 AI Explainability: Commit to enhancing the intelligibility of AI decisions (e.g., product recommendation logic) for stakeholders such as franchisees and consumers to avoid "black-box" operations that damage trust.

6.4 Appeals & Human Review: Stakeholders have the right to appeal; the company provides transparent redress procedures and human-led review pathways.