

COUNTRY GARDEN SERVICES HOLDINGS COMPANY LIMITED

碧桂園服務控股有限公司

(Incorporated in the Cayman Islands with limited liability)

(Stock Code: 6098.HK)

ARTIFICIAL INTELLIGENCE AND AUTOMATED DECISION-MAKING MANAGEMENT POLICY

1. General Provisions

1.1 Country Garden Services Holdings Company Limited (hereinafter referred to as "Country Garden Services" or the "Group"), in order to regulate the full lifecycle management of the Group's artificial intelligence and automated decision-making tools, implement the nine core requirements of data privacy protection, cybersecurity protection, algorithm fairness, human-in-the-loop review, explainability, clear division of responsibilities, capability boundary control, ecological footprint reduction, and prohibition of non-compliant AI systems, and prevent risks such as algorithmic bias, privacy breaches, cyber intrusions, loss of decision-making control, illegal surveillance, manipulative practices, excessive consumption of computing resources, etc., hereby formulates this Policy in accordance with the Cybersecurity Law, the Data Security Law, the Personal Information Protection Law, the Interim Measures for the Management of Generative Artificial Intelligence Services, the Provisions on the Security Management of Facial Recognition Technology Application, the classified cybersecurity protection standards, internal control requirements for listed companies, and the compliance principles of the EU Artificial Intelligence Act, so as to safeguard the legitimate rights and interests of property owners and employees, as well as the Group's operational security.

1.2 The Group's highest AI review and decision-making body is the Group's Executive Management. The launch of high-risk and major new AI systems must be reviewed and approved by the Executive Management.

2. Scope of Application

The Group and all its subsidiaries, regional project branches, and all employees of the Group's entire operations, must comply with this Policy. It covers all self-developed and externally purchased AI models and automated decision-making systems, encompassing the entire process from requirements design, development and training, launch and operation, iterative updates, to decommissioning and destruction. The deployment, access, or use of AI systems that fall under the prohibited categories of the EU Artificial Intelligence Act is strictly forbidden.

3. Terms and Definitions

3.1 Artificial Intelligence System: Software and hardware that rely on technologies such as machine learning, computer vision, and large models to automatically complete data identification, analysis, and judgment, including self-developed AI and third-party AI.

3.2 Automated Decision-making: AI that autonomously outputs business conclusions and automatically executes actions, e.g., owner access control, fee collection, facial recognition access, equipment dispatching, etc.

3.3 Human-in-the-Loop: A mechanism whereby high-risk AI processes are equipped with manual review checkpoints; AI provides only reference conclusions, and key operations must be confirmed and executed by human personnel.

3.4 Algorithmic Bias: Discriminatory or unfair outcomes produced by a model against different owners or groups due to defects in training data or design.

3.5 AI Explainability: AI decisions are supported by complete traceable evidence and can clearly articulate the logic behind the determination.

3.6 AI Capability Boundary: A clearly defined scope of business and data that the AI is permitted to handle, along with a list of prohibited operations.

3.7 Ecological Footprint: The comprehensive environmental impact of AI models and supporting computing data centers throughout their full lifecycle, including electricity consumption, water consumption, hardware resources, carbon emissions, electronic waste, and other environmental effects; the overall environmental load shall be minimized.

3.8 Non-compliant AI System: AI systems engaged in manipulative practices, exploitation of system vulnerabilities, social scoring, unauthorized biometric surveillance, and any other AI systems explicitly prohibited by the EU Artificial Intelligence Act.

3.9 AI Security Incident: Incidents including privacy breaches, cyber intrusions, algorithmic discrimination, loss of decision-making control, use of AI models beyond the authorized scope, operation with excessive ecological footprint, etc.

4. Core Management Principles

4.1 Privacy-First Principle: Data collection shall follow the principle of minimum necessity. Personal and biometric sensitive information must be strictly protected; unauthorized or out-of-scope use of data is prohibited.

4.2 Cybersecurity Principle: AI models, data storage, and call chains must be protected by network security measures throughout the entire process, blocking intrusion and data exfiltration channels, and strictly controlling the risks of internal-external network data exchange.

4.3 Fairness and Non-Discrimination Principle: Training data and algorithms shall be validated to eliminate discriminatory biases, ensuring fair and just algorithmic determinations.

4.4 Human-in-the-Loop Principle: High-risk business processes shall retain full human intervention authority; AI shall not be permitted to fully autonomously execute critical business operations.

4.5 Transparency and Explainability Principle: The reasoning chain of AI decisions shall be retained to ensure traceability and explainability; black-box determinations are prohibited.

4.6 Transparency and Identifiability Principle: AI-generated content and decision outputs shall be distinctly labeled with both explicit and implicit labels to ensure users can recognize the AI nature; the first interaction in any AI-powered service must inform the user accordingly.

4.7 Clear Accountability Principle: Accountability follows the rules of "who uses, who is responsible" and "who develops, who bears liability"; the responsible party for AI-generated outcomes must be clearly defined.

4.8 Boundary Control Principle: The AI capability boundary must be defined prior to launch; operation of models beyond the prescribed business scope is strictly prohibited.

4.9 Green and Low-Carbon Principle: Lightweight models and green computing architectures shall be selected; computing scheduling shall be optimized to reduce the overall AI ecological footprint and minimize environmental burden.

4.10 Compliance and Prohibition Principle: The deployment or activation of AI systems engaged in manipulative practices, vulnerability exploitation, social scoring, unauthorized biometric surveillance, or other legally prohibited AI systems is strictly forbidden.

5. Responsibilities and Division of Duties

5.1 Business User Departments: Coordinate AI business requirements; define AI capability boundaries; conduct ethical training on AI use and AI result review; establish an appeals process for users to contest AI decisions and complete manual re-review; promptly report security risks and abnormal high-energy-consumption operations; assume primary business responsibility for AI decisions in their respective scenarios.

5.2 Digital Capability Center (AI Development and Implementation Team): Responsible for technical management of self-developed AI; implement data masking and de-identification and privacy protection, network security, algorithmic bias correction, decision traceability; select lightweight models to reduce ecological footprint; conduct compliance checks to eliminate non-compliant AI models; complete risk assessments for launch and iterative updates; retain full-process audit logs; ensure secure destruction of data upon decommissioning; bear technical responsibility for the models. Launch plans for major high-risk AI systems must be submitted to the Group's Executive Management for approval.

5.3 Information Security Team: Oversee AI compliance management and external model access control; regularly conduct security, algorithm, and computing energy consumption audits; verify ecological footprint management; handle AI security incidents; coordinate with classified protection (MLPS) and third-party audit work.

6. Full-Process Management Requirements for AI

6.1 Launch Admission and Boundary Control: Prior to launch, complete a comprehensive compliance risk assessment and retain records; define capability boundaries and computing energy consumption management plans; verify that the system is not a non-compliant AI system. High-sensitivity scenarios such as facial recognition and finance, as well as major new AI systems, must be reviewed and approved by the Group's Executive Management. The system shall not overstep its authority to perform high-risk operations such as approvals, fund deductions, or permission grants. Overall ecological footprint shall be strictly controlled. Deployment of manipulative practices, vulnerability exploitation, social scoring, unauthorized biometric surveillance, and other prohibited AI systems is strictly forbidden.

6.2 Data Privacy Protection: For sensitive information such as facial data, separate written authorization must be obtained; data shall be stored in isolated internal network environments and retained for the shortest necessary period. Training data must be appropriately masked or de-identified. Transmission of sensitive data to external AI systems is strictly prohibited; internal data export must follow approval procedures; highly sensitive data shall generally not be exported or downloaded.

6.3 Cybersecurity Protection: AI databases and models shall be deployed in isolated security zones, with firewalls, access controls, and interface authentication configured; regular vulnerability scanning and patching shall be conducted. AI may only be used through compliant internal network channels; unauthorized transmission of sensitive data to external models is prohibited.

6.4 Algorithmic Bias Prevention: Training samples shall be balanced and discriminatory features removed; regular fairness testing of algorithms shall be conducted; models found to have bias shall be corrected promptly to eliminate unfair discriminatory outcomes.

6.5 Human-in-the-Loop Control: In high-risk scenarios such as tenant access, large-value fee adjustments, bulk payment collection actions, and facial recognition permission granting, AI shall provide only recommendations, and manual review and confirmation shall be mandatory before execution. The system shall allow manual override and termination of automated AI processes at any time. Low-risk scenarios shall be subject to periodic sample review.

6.6 Explainability Management: AI decisions shall retain complete traceability logs to enable explanation of the basis for determinations to owners and auditors. For external services, explanations shall be provided in easy-to-understand language.

6.7 Operation and Maintenance & Decommissioning Management: Regularly conduct inspections for vulnerabilities, algorithmic deviations, and computing energy consumption. For major version iterations, a new compliance risk assessment must be completed. Continuously optimize model architectures to reduce ecological footprint. Upon decommissioning, training data and models shall be thoroughly destroyed. For third-party models upon contract expiration, confirm data deletion. Sensitive AI output documents shall be watermarked and their transmission channels strictly controlled.

7. AI Security Incident Handling and Auditing

7.1 Graded Handling of AI Security Incidents

Incident Level	Typical Scenarios	Handling Timeframe	Handling Requirements
Level 1 Major Incident	Bulk facial data leakage, significant financial loss, model theft, large-scale algorithmic discrimination, prolonged excessive ecological footprint operation	Within 1 hour	Shut down the involved AI; report to management; conduct source tracing and investigation; notify affected owners; retain audit logs and report to authorities as required.
Level 2 General Incident	Localized algorithmic bias, unauthorized external AI access, third-party model ecological footprint exceeding limits	Within 24 hours	Implement comprehensive rectification; restrict relevant personnel permissions; conduct compliance training; optimize computing scheduling.
Level 3 Minor Incident	Minor overstepping of boundaries, unauthorized calling of external AI without sensitive data, short-term abnormal energy consumption	Same day	Complete rectification; issue internal notice; improve boundary controls and lightweight optimization plans.

7.2 Regular Audits: The Information Security Team shall conduct regular AI special audits, verifying the implementation of privacy protection, cybersecurity, algorithmic fairness, human-in-the-loop, explainability, capability boundaries, and ecological footprint management, forming a closed-loop for rectification. Complete assessment reports, audit logs and remediation records shall be retained for annual MLPS assessments and third-party audits.

7.3 Full-process AI audit logs shall be retained for no less than 3 years to meet compliance and traceability requirements.

8. Policy Review and Amendment

This Policy has been reviewed and approved by the Board of Directors of the Group. The Executive Management and the Information Security Team are responsible for overseeing its effective implementation. The applicability and effectiveness of this Policy shall be reviewed and evaluated from time to time. Amendments to this Policy shall be considered based on review results, updates to laws and regulations, developments in industry events, and the Company's actual operational circumstances, so as to continuously enhance the Company's sustainable development performance.