

ISOLVED

ARTIFICIAL INTELLIGENCE USE, TRANSPARENCY, AND ACCEPTABLE USE POLICY

Effective Date: August 1, 2026

ISOLVED (“Company,” “we,” “our,” or “us”) uses artificial intelligence (“AI”) in our products and services to enhance functionality, improve efficiency, and deliver meaningful value to our customers. We believe AI should be deployed in a manner that is responsible, transparent, and aligned with our customers’ expectations for security, reliability, and control.

While AI can improve outcomes, it also introduces new and evolving risks, including risks related to accuracy, bias, security, and unintended outputs. We have implemented technical, contractual, and organizational controls to manage these risks as part of our broader trust, security, and governance framework.

This Policy explains: (a) how we use AI in our products and services; (b) how we protect customer data in connection with AI; and (c) the rules governing the use of AI-enabled features. It is intended to provide customers with clear, accurate, and practical information about how AI operates within our products and the safeguards that apply.

Unless otherwise expressly stated in applicable agreements, customers remain responsible for evaluating and validating AI-generated outputs before relying on them, particularly in connection with business-critical, legal, financial, or operational decisions.

1. OUR APPROACH TO TRUSTED AI

1.1 Design and Deployment. We design and deploy AI in accordance with the following principles:

- (a) **Transparency:** Customers are informed when AI is used in a material way.
- (b) **Human Accountability:** AI augments and/or offers suggestions, but does not replace, human judgment. Usage for decisions is overseen and affirmed by human decision makers, be it you or Isolved personnel.
- (c) **Privacy and Security:** Customer data is protected through strict technical and contractual safeguards.
- (d) **Fairness and Non-Discrimination:** AI systems are designed and monitored to reduce bias and harmful outcomes.
- (e) **Reliability and Safety:** AI outputs are tested, monitored, and continuously improved.
- (f) **Governance:** AI use is subject to cross-functional governance, risk-based approvals, and periodic review.
- (g) **Risk-Based Approach:** AI systems are classified and governed based on their potential impact and risk profile.
- (h) **Security and Privacy by Design.** AI systems are developed with embedded security and privacy safeguards throughout their lifecycle.
- (i) **External validation and Standards alignment.** We align, where applicable, with industry best practices and applicable laws, including but not limited to data privacy laws.

These principles are implemented through internal controls and oversight designed to ensure that AI systems are appropriate for their intended use and risk profile. For use cases involving regulated activities (e.g., financial services, healthcare, employment-related decisions), additional safeguards may apply, including enhanced testing, review, and human oversight consistent with applicable regulatory requirements.

1.2. How We Use & Incorporate AI In Our Products and Services. We primarily use AI for assistive, augmentative, or suggestion purposes, designed to enhance, not replace, human decision-making. AI is not used for fully autonomous decision-making that produces legal, financial, employment, or similarly significant effects unless explicitly disclosed and subject to appropriate safeguards. Customers remain solely responsible for employment, payroll, tax, compensation, benefits, hiring, promotion, disciplinary, scheduling, termination, and other workforce-related decisions, regardless of whether AI-assisted functionality is used in connection with such decisions. Examples of our use cases include:

- (a) content generation, summarization, and recommendations;
- (b) data analysis, forecasting, and insights;
- (c) workflow automation and decision support;
- (d) customer support and conversational interfaces; and
- (e) other domain-specific capabilities as described in applicable product documentation.

AI may generate outputs, recommend outcomes or actions, prioritize or rank information, assist with decision-making, or automate certain processes depending on the feature or functionality.

1.3 Monitoring and Continuous Improvement. We monitor AI systems for accuracy and performance, bias and unintended outcomes, incident response triggers, and misuse. We maintain processes to detect, respond to, and remediate AI-related incidents, including unintended outputs, bias, or security vulnerabilities. We use customer feedback, human review cycles, testing, and monitoring to improve AI systems over time. Customers may report concerns, errors, or unexpected outputs through our support channels, which are incorporated into our monitoring and improvement processes.

1.4 We Are Committed to Transparency & Disclosure

- (a) We provide clear and conspicuous disclosure when users interact with AI systems where such use is not reasonably apparent.
- (b) AI-generated outputs may be labeled, annotated, or otherwise identified, including through user interface indicators or contextual disclosures, where appropriate.
- (c) We do not misrepresent AI-generated content as human-generated.
- (d) Representations regarding AI capabilities are accurate, substantiated, and not misleading, consistent with regulatory expectations.
- (e) Where appropriate, we provide documentation or user-facing explanations describing AI functionality and limitations.

1.5 AI Limitations. Innovation concerning AI is occurring quickly and providing significant benefits to our customers.

- (a) Although AI can improve productivity and insights, it does not replace human judgment. All AI has limitations and hallucination risks, and its accuracy and reliability should be considered accordingly given that AI outputs inaccurate. Incomplete, inconsistent, outdated, or fabricated. AI outputs are probabilistic in

nature and performance may vary over time due to changes in models, underlying data, and operating conditions . Further, AI-generated content should not be relied upon without appropriate review, particularly for business-critical, legal, financial, or operational decisions.

- (b) For higher-risk use cases of AI in our product, we implement enhanced controls, including:
 - i. testing and validation;
 - ii. monitoring and performance review; and
 - iii. human oversight where appropriate.
- (c) We do not guarantee that AI outputs are error-free or suitable for all use cases.

When used appropriately and with human oversight, AI can enhance productivity, insight generation, and decision support. Our approach is designed to enable responsible innovation where we leverage the benefits of AI while implementing enterprise-grade safeguards that support trust, reliability, and accountability.

1.6 AI outputs are provided for informational and assistive purposes only and are not intended to constitute legal, financial, or professional advice.

2. HUMAN OVERSIGHT AND CUSTOMER CONTROL

2.1 We design AI systems to include appropriate human-in-the-loop mechanisms, particularly for higher-risk use cases. Where applicable, customers may:

- (a) review, edit, or override AI-generated outputs;
- (b) request human assistance; or
- (c) escalate or challenge AI-generated outcomes.

Certain safeguards may be enabled by default, while others may be configurable depending on the specific feature and product design.

2.2 AI systems are not intended to replace human decision-making in contexts involving material risk (financial, legal, employment, pricing, and other regulated data). Appropriate human review and validation should be performed before action is taken.

2.3 Customers are responsible for determining whether to utilize AI imbedded Services or otherwise enable AI features within their environment, and review processes and access controls consistent with their internal policies and risk tolerance. For AI imbedded services or features that provide direct control of enablement, access controls, or configuring settings to Customer, Customer is solely responsible for whether and how they set up their environment.

3. EXPLAINABILITY AND ACCOUNTABILITY

3.1 Where technically feasible, we maintain documentation and governance processes that support traceability and auditability of AI systems.

3.2 Where appropriate or required by law, we provide customers with meaningful information regarding:

- (a) the general functionality and intended use of AI-enabled features
- (b) how AI-generated outputs are produced;
- (b) key factors influencing outputs;

- (c) known limitations of the system; and
- (d) applicable user responsibilities.

3.3. At an appropriate level of abstraction, we maintain internal documentation describing AI system design, intended use limitations, and evaluation metrics.

3.4 We implement controls to support appropriate oversight and auditability of AI system behavior in a manner consistent with data minimization and security requirements.

3.5 Subject to applicable agreements, we may provide customers with information reasonably necessary to support their compliance, audit, or risk assessment requirements related to AI features.

4. CUSTOMER DATA, PRIVACY, AND MODEL TRAINING

4.1 We may process customer data as inputs to AI systems (“inference”) to generate outputs requested by the customer. We do not use customer data, inputs, or outputs to train or fine-tune AI models unless explicitly permitted by our contract with the customer. Any permitted use for model improvement will be subject to applicable contractual, privacy, security, and data protection obligations. Where applicable, we apply controls to limit the reuse of customer-derived data and outputs for model improvement, including technical and contractual restrictions designed to prevent unauthorized use.

4.2 We apply data minimization principles and limit data use to what is necessary for the functionality provided.

4.3 Customers are responsible for ensuring that data they input complies with applicable laws and their own internal policies.

4.4 Unless otherwise explicitly stated, we do not sell customer data, including data processed in connection with AI features. We do not disclose customer data to third parties for their independent marketing, advertising, or data monetization purposes. Any sharing of data with service providers or vendors is performed solely to deliver and support our products and services, subject to contractual data protection obligations.

5. VENDORS AND AI TECHNOLOGY.

5.1 We may use third-party AI technologies integrated with our Service for your use or as part of our own Services. AI vendors are subject to internal security, privacy, and compliance review appropriate to the role they play in delivering the service.

5.2 We also implement contractual safeguards to ensure vendors do not use customer data for unauthorized purposes, provide appropriate data protection commitments, and comply with applicable laws.

5.3 We will evaluate material changes to vendor AI technologies, including changes to models, features, or data handling practices, to ensure continued alignment with our security, privacy, and compliance requirements.

6. PRODUCT INTEGRATION AND CUSTOMER-FACING AI

6.1 AI features embedded in our products are designed to be transparent, reliable, and subject to appropriate controls. Where AI outputs may materially impact customers, outputs are subject to validation or oversight and human review mechanisms are available.

6.2 Certain AI features may be optional or configurable, while others may be embedded within core product functionality. Where feasible, we provide customers with the ability to enable, disable, or configure AI-driven features.

6.3 Customer agreements, product documentation, and disclosures (e.g. labeling, user interface indicators, disclaimers at point of use) accurately reflect the role of AI in the product, limitations of AI outputs, and our data usage practices.

6.4 AI capabilities within our products may include different types of systems, such as generative AI (e.g., content generation), predictive or analytical models (e.g., forecasting or scoring), and automation tools (e.g., workflow assistance). These systems may operate differently and present different risk profiles. We apply controls appropriate to the nature, purpose, and impact of each type of AI functionality.

7. DATA HANDLING, RETENTION, AND SECURITY

7.1 AI-related data is retained only as necessary for functionality, compliance, and improvement in accordance with applicable retention schedules, operational requirements, legal obligations, contractual obligations, and legitimate business purposes.

7.2 We implement controls that may include: data retention and deletion, anonymization and aggregation, and monitoring appropriate to the feature and risk profile.

7.3 We maintain safeguards to prevent unauthorized access, misuse, or disclosure of data.

7.4 We ensure that our AI-related practices are aligned with and governed by our broader information security and data protection programs, which may include industry-standard frameworks (e.g., SOC 2) where applicable.

7.5 We maintain processes to detect, investigate, and respond to AI-related incidents, including those involving data security, unintended outputs, or system misuse. Where required by applicable law or contract, we will notify affected customers of material incidents involving their data or use of AI features, consistent with our incident response and data breach notification obligations.

8. CUSTOMER ACCEPTABLE USE OF AI FEATURES

8.1 Customers are responsible for using AI-enabled features in a lawful and responsible manner. Misuse may expose you to legal/regulatory risk and/or may impact system integrity.

8.2 Customers may not use AI features to:

- (a) engage in unlawful, harmful, or fraudulent activities;
- (b) generate or distribute content that is:

- i. abusive, harassing, or discriminatory;
 - ii. materially misleading or deceptive;
 - iii. infringing on intellectual property rights;
- (c) develop or support:
 - i. harmful activities;
 - ii. exploitation or abuse of individuals;
 - iii. unauthorized surveillance or biometric identification;
- (d) create or disseminate:
 - i. deepfakes or manipulated content intended to mislead;
 - ii. profiling or targeting based on protected characteristics;
- (e) make or automate decisions that have legal, financial, employment, pricing, or similarly significant effects without appropriate human oversight;
- (f) attempt to reverse engineer, extract, or misuse AI models or underlying systems;
- (g) attempt to exploit vulnerabilities in AI systems, including prompt injection or adversarial inputs;
- (h) input or process data in violation of applicable laws or contractual obligations; or
- (i) use AI outputs in a manner that violates applicable laws or regulations.

9. ENFORCEMENT AND VIOLATIONS

9.1 Violations of this Policy may result in:

- (a) suspension or termination of access to AI features;
- (b) restriction of functionality; or
- (c) other actions consistent with applicable agreements.

9.2 We reserve the right to investigate suspected violations and take reasonable and proportionate action to protect our systems, customers, and users. Customers agree to cooperate in reasonable investigations of violations.

9.3 Violations may constitute a breach of applicable customer agreements.

10. POLICY UPDATES & CONTACT INFORMATION

10.1 AI is subject to evolving legal and regulatory requirements. We monitor applicable laws, regulations, and industry standards and will update our practices and this Policy as necessary to maintain compliance and reflect emerging expectations.

10.2 We may update this Policy periodically to reflect changes in technology, applicable law, or our practices. Where required by law or where changes are material, we will provide notice to customers in accordance with applicable agreements. Updates will be posted on our website and become effective as specified in the notice.

10.3 For questions regarding this Policy or our use of AI, contact us at Legal@isolvedhcm.com.