

## Acceptable AI Policy (AAIP)

*Governing the permitted and prohibited uses of the Pathify AI Agent. Organized by Agent function and by how each function is treated under Colorado Senate Bill 26-189 and analogous frameworks.*

### Executive Summary

Pathify's AI Agent helps users find information, navigate institutional resources, and reach the right people faster. This Policy explains, in plain terms, what the AI Agent does, what it is not allowed to do, how user information is handled, and what safeguards apply to each function. It is organized by the Agent's functions rather than by use tiers, because the functions differ in how the law treats them, and a function-based description is both more accurate and easier for an institution, a user, or a regulator to follow. The AI Agent serves multiple End User types, including students, prospects, applicants, alumni, parents or guardians where authorized, and staff; where an End User is an enrolled student, that user's institutional records are FERPA (Family Educational Rights and Privacy Act) education records, handled as described in Section 4.3.

**This Policy applies the law to each Agent function individually and states the result openly.** Under Colorado Senate Bill 26-189 ("SB 26-189") and similar laws, the question is not whether a system is "AI" but whether it materially influences a consequential decision about a person. Pathify takes the position that the Agent's core question-and-answer function does not constitute covered automated decision-making technology ("Covered ADMT"). Safety, and when released, Monitoring detect a possible concern and escalate it to a qualified institutional staff member, who decides whether to act. Because the human, not the Agent, makes any consequential decision, neither Safety nor Monitoring is Covered ADMT. As of the date of this Policy, no Agent function is Covered ADMT. Treating each function according to how the law applies to it is the approach this Policy enforces.

**A human, not the Agent, makes every consequential decision.** The Agent informs, retrieves, classifies, and refers. Where a function bears on a decision that affects a student's educational path, a qualified institutional staff member makes the actual determination, with the Agent's output as input and not as the decision. This Policy calls that an Independent Determination, and it is the single most important safeguard across SB 26-189, the California Consumer Privacy Act automated-decisionmaking regulations, the European Union AI Act, Article 22 of the European Union General Data Protection Regulation, the parallel Articles 22A to 22C of the United Kingdom General Data Protection Regulation, and other similar laws and regulations.

**This Policy helps your institution stand up and document its own AI governance framework.**

**Record information can reach the model provider.** When an institution configures the Agent to use an individual's own record context (for example, grades, enrollment status, or financial aid information) to answer a question or perform a function, those record values are transmitted

to the model provider (Google) to generate the response. This Policy describes that data flow accurately and applies the Family Educational Rights and Privacy Act (“FERPA”) school official framework to student records within that flow. It does not represent that record values are withheld from the model provider.

**The Policy is principles-based and built for procurement review.** It is designed to work across the principal AI governance regimes the institutional customer base intersects, including the United States National Institute of Standards and Technology (“NIST”) AI Risk Management Framework, SB 26-189 and other emerging United States state AI laws, the New York public-sector AI governance expectations, the European Union AI Act, the United Kingdom Data (Use and Access) Act 2025 and Information Commissioner’s Office guidance, the Australian Privacy Act 1988 (Cth), the New Zealand Privacy Act 2020, ISO/IEC 42001, and the OECD AI Principles. The underlying protections operate independent of any single statute’s outcome.

**What this Policy does, in five points.**

- Describes each Agent function in plain terms and states, function by function, whether it is treated as Covered ADMT.
- Keeps a qualified human in control of every consequential decision, enforced by both policy and system architecture.
- Describes the data flow to the model provider accurately, including that record values can be transmitted when the institution so configures, and applies the FERPA school official framework to those transmissions.
- Addresses unlawful-discrimination risk through testing, monitoring, and documentation that would apply to any function ever classified as Covered ADMT.
- Allocates responsibilities between Pathify (developer) and the institution (deployer), with FERPA processes used where legally available to support deployer notice and consumer-rights obligations for education decisions.

## **1. Purpose and Scope**

This Acceptable AI Policy (this “Policy” or “AAIP”) governs the permitted and prohibited uses of the artificial intelligence features (“AI Features”) included within the services Pathify provides under its Master Services & Software License Agreement (“MSA”). The principal AI Feature is the Pathify AI Agent (the “AI Agent” or “Agent”), a conversational interface powered by Google Gemini technology that retrieves institutional information, answers questions, and performs the functions described in Section 2 in response to user prompts.

This Policy serves a dual purpose. First, it establishes the usage rules and safeguards that govern how the AI Agent operates as an institutional tool under customer direction and control. Second, it documents Pathify’s assessment of each AI Agent function under SB 26-189 (codified within Colorado Revised Statutes (“C.R.S.”) Title 6, Article 1, Part 17, as enacted and effective January 1, 2027) and provides a framework that customers may use to support their own AI governance, notice, and impact-assessment processes. SB 26-189, together with similar state automated-decision-making laws, is referred to as the “Applicable AI Laws.”

Pathify is adopting this Policy in advance of the January 1, 2027 effective date of SB 26-189 to establish its contract posture and operating practices ahead of that date. References to SB 26-189 requirements describe obligations that apply to consequential decisions made on or after January 1, 2027. This Policy applies to Pathify, its employees and contractors, and to all institutional customers (“Customers”) and their end users (“End Users”) who interact with the AI Features.

### **1.1 Intended Use and Statutory Positioning**

The AI Agent is intended to provide information, navigation, answers to questions, informational referrals, and informational suggestions about institutional resources, contacts, and processes, and to generate related content, in each case to assist End Users and institutional staff. The AI Agent’s informational suggestions inform an End User’s or staff member’s own decision and are not recommendations used to make or materially influence a Consequential Decision. The AI Agent is not intended, and this Policy prohibits its use, to make or materially influence a Consequential Decision through its core question-and-answer function.

Pathify does not take the position that every AI Agent function is outside the scope of the Applicable AI Laws. Pathify’s positions, stated function by function in Section 2 and carried through the rest of this Policy, are: the core question-and-answer function does not materially influence a Consequential Decision and is not Covered ADMT; the Safety and Monitoring functions are detection-and-escalation functions where an Institution Designee, not the Agent, makes any consequential decision, and are not Covered ADMT. This Policy is structured so that a regulator, institutional leader, or procurement reviewer can see exactly which functions are treated as regulated and how the corresponding obligations are met.

Where any function approaches the line between assistance and a Consequential Decision, this Policy requires that a qualified Institution Designee exercise an Independent Determination between the AI Agent’s output and any institutional decision. This human-decision-maker requirement supports two distinct analyses without conflating them. For SB 26-189, a genuine Independent Determination is intended to keep the AI Agent from being a non-de-minimis factor that materially influences a Consequential Decision. For FERPA, the school official exception under 34 C.F.R. 99.31(a)(1)(i) does not turn on who approves a downstream outcome; it requires that the institution maintain direct control over the outside party’s use and maintenance of education records, that use be limited to the purpose of the disclosure, that redisclosure be restricted, and that the outside party meet the criteria in the institution’s annual FERPA notification. Those FERPA conditions are addressed in Section 4.3 and depend on the institution-Pathify and Pathify-Google arrangements and configuration, not on the Independent Determination alone. The two regimes align in that the institution remains the decision-maker and the controller of its records, but they are not coextensive, and satisfaction of one is not represented as satisfaction of the other.

### **1.2 Defined Terms**

Capitalized terms used in this Policy and not defined in the MSA or herein have the meanings set forth below. To the extent of any conflict between a definition in this Policy and a definition in

the MSA, the MSA controls. Capitalized terms used without definition have the meanings given in SB 26-189.

- **“Institution-Approved Content”** means content the Customer designates as available for retrieval and surfacing by the AI Agent, consisting, as and to the extent the Customer determines, of (i) content published on the Customer’s institutional websites, (ii) Pathify-provided white-labelled content approved by the Customer, and (iii) third-party web content specifically approved by the Customer for AI Agent retrieval.
- **“Institution Designee”** means a natural person employed by, contracted to, or otherwise affiliated with the Customer and designated by the Customer as a school official under FERPA (34 C.F.R. 99.31(a)(1)) to receive AI-Agent-originated referrals, classifications, or evaluations, with a defined scope of legitimate educational interest in those communications.
- **“Core Question-and-Answer Function”** means the AI Agent function that retrieves and surfaces Institution-Approved Content and an individual’s own institutional records in response to End User prompts, answers questions, and makes informational referrals, including where the institution configures the Agent to use the individual’s own record context to generate a response.
- **“Generative Fallback”** means the AI Agent’s generation of a response from the underlying model’s general capabilities where a prompt is not answered from Institution-Approved Content or institutional records, the substance of which is not bounded by institution-approved sources.
- **“Safety Function”** means the AI-based classification function that screens conversation content to identify a possible risk of harm (for example, expressions of self-harm or threats of violence) and, where the Customer has configured it, routes a notification to a designated Institution Designee for that Designee’s independent judgment and response.
- **“Monitoring Function”** means the AI-based classification function that, where the Customer has enabled it, identifies conversations indicating institution-designated concerns (for example, academic disengagement or a need for additional support) and routes a notification to a designated Institution Designee for that Designee’s independent judgment and response.
- **“Independent Determination”** means a determination that an Institution Designee makes based on the Designee’s own review of the relevant records and criteria and the Designee’s own judgment. Where the AI Agent surfaces information to a Designee, for example a Safety or Monitoring escalation, the Designee decides whether and how to act on that information; the Designee is not constrained by the AI Agent output, may decline to act on it, and may reach a conclusion that exceeds or differs from it. An Independent Determination is not a mere approval or rubber-stamp. To qualify, (i) the Designee has access to the underlying records necessary to evaluate the matter, not merely the AI Agent output; (ii) the Designee has, and may exercise, genuine authority to decide the matter independently of the AI Agent output; (iii) the Designee records the basis for the determination sufficient to show independent judgment; and (iv) the Customer maintains review controls reasonably designed to confirm determinations are not made by default

acceptance, including periodic quality-assurance sampling. Where these conditions are not met for a workflow such that the AI Agent output operates as a non-de-minimis factor in a Consequential Decision, the AI Agent may be materially influencing that decision, and the contingent Covered ADMT procedures in Sections 8.1 and 9.1 apply to that workflow. The Customer is responsible for implementing, in its environment, the staffing, review authority, and recordkeeping that make each Independent Determination genuine; Pathify states this requirement and provides supporting product capabilities but does not control or guarantee the institution's internal decision-making.

- **“Covered ADMT”** means automated decision-making technology, as defined in SB 26-189, that is used to materially influence a Consequential Decision. This Policy takes the position that no current AI Agent function is Covered ADMT. The Core Question-and-Answer Function is informational; the Safety and Monitoring functions detect and escalate to an Institution Designee, who makes any consequential decision. If an AI Agent function were ever used to materially influence a Consequential Decision, this Policy would treat that function as Covered ADMT, and the contingent procedures in Sections 6.1, 7, 8, 8.1, 9.1, and 9.2 would apply to it.
- **“Consequential Decision”** means a decision that has a material effect on, or that determines, a consumer's access to, eligibility for, or the cost or terms of, education enrollment or an education opportunity, or any other decision identified as a consequential decision under SB 26-189. To the extent of any conflict with the MSA, the MSA controls.
- **“Materially Influence”** means to be a non-de-minimis factor in making a Consequential Decision, including by constraining, ranking, scoring, recommending, classifying, or otherwise meaningfully altering how a Consequential Decision is made, consistent with SB 26-189. “Materially Influences” has a corresponding meaning.
- **“ADMT Information-Tool Exclusion”** means the exclusion from the definition of automated decision-making technology under C.R.S. 6-1-1701(2)(b)(III) for technology that communicates with consumers in natural language to provide information, make referrals or recommendations, answer questions, or generate other content, where the technology is not contracted, advertised, marketed, configured, or intended to be used in a Consequential Decision and is subject to an acceptable use policy that prohibits generated content from being used in a Consequential Decision. This Policy relies on the exclusion only as a secondary and additional basis for the Core Question-and-Answer Function, and does not rely on it for the Monitoring Function, which is outside Covered ADMT because an Institution Designee, not the Agent, makes any consequential decision.
- **“NIST AI RMF”** means the National Institute of Standards and Technology AI Risk Management Framework 1.0.
- **“Integrated Action”** means the AI Agent's triggering of a deterministic action in the Customer's or a connected institutional system from within the AI Agent user experience (also called write-back), as described in Section 3.5. An Integrated Action is available only where the Customer has first configured or built it, and is subject to the four conditions in Section 3.5.

### 1.3 Procurement and Governance Use of this Policy

This Policy is intended to serve both as an operational acceptable-use policy and as part of Pathify's AI governance and procurement-readiness documentation for higher education Customers. It is designed to support institution-side review by procurement, privacy, security, audit, academic leadership, and public-sector oversight functions that require a clear statement of AI feature scope, data handling, human review, and accountability. It should be read together with the MSA, the Data Processing Addendum ("DPA"), and Pathify's supporting governance materials: the MSA allocates commercial authority and remedies, the DPA governs processing and privacy and security commitments, this Policy governs AI feature use and function-specific safeguards, and supporting materials may provide additional product-surface and data-flow detail.

As an institution-facing governance resource, this Policy is structured to help a Customer support its own AI-framework obligations. For public higher-education systems, whose systemwide policies increasingly require each campus to adopt local AI policies addressing data protection, bias evaluation, preserved human decision-making, role clarity, procurement safeguards, and training, this Policy provides vendor-side documentation a campus can use and cite, summarized in the Public Higher-Education AI Framework Crosswalk in Section 10.8. It is also written to address AI requirements in other states where Customers operate today (Section 10.1) and to be extended as additional state and federal AI requirements take effect, through the framework-alignment, crosswalk, and review mechanisms in Sections 10 and 12. This Policy does not warrant that any institution's obligations are fully satisfied by Pathify's materials alone; each Customer retains primary responsibility for meeting its own obligations, as provided in Section 10.5 (No Warranty of Sufficiency).

## **2. AI Agent Functions and Classification**

The AI Agent performs the functions described in this Section 2. Each function is described in plain terms and then classified under SB 26-189. The functions are not "tiers" and are not sequenced; an institution may have several enabled at once, subject to configuration and the safeguards in this Policy. The classification of each function is summarized in Section 2.5 and carried through Sections 3 through 10.

### **2.1 Core Question-and-Answer Function**

The Core Question-and-Answer Function is the AI Agent's primary operation. The Agent searches Institution-Approved Content and other locally available institutional information, and where the institution so configures, an individual's own institutional records, and responds to End User prompts with text and data displays. It includes answering general campus questions (locations, hours, events, procedures, deadlines); providing navigation within Pathify's Campus Experience Platform ("CXP"); routing users to the correct office, department, or contact; surfacing institutional frequently asked questions, directory listings, and published content; retrieving and displaying institution-specific or user-specific records (for example, grade point average, enrollment status, or academic history) from platform data sources; and performing web searches when the Customer enables that option.



**Configurable record context.** Whether the Agent uses an individual's own record values to answer a question is determined by the institution's configuration. When the institution enables the use of the individual's record context to generate a response, the record values themselves, and not merely a reference to them, are transmitted to the model provider as needed for that purpose, as described in Section 2.6 and Sections 4.1 and 4.3. This Policy does not represent that record values are withheld from the model provider. Even where the Agent uses an individual's own record values to answer, the answer is informational: the Agent presents or explains existing institutional records and does not evaluate the individual, score or rank the individual, or generate an output used to make or guide a Consequential Decision. Any Consequential Decision remains with the institution and its qualified staff.

**Classification: not Covered ADMT.** The Core Question-and-Answer Function presents records and routes inquiries; it does not evaluate a student against criteria or generate an output used to make or guide an institutional determination. It is therefore not a non-de-minimis factor in any Consequential Decision and is not Covered ADMT under SB 26-189. The Core Question-and-Answer Function also falls within SB 26-189's exclusion from the definition of Consequential Decision for routine academic administration and student-support processes that do not materially influence a Consequential Decision (C.R.S. 6-1-1701(3)(b)(IX)). As a secondary and additional basis, and only for this function, the configuration and operation described here, together with the prohibition in Section 3.1, are intended to satisfy the ADMT Information-Tool Exclusion. This Policy does not rely on that exclusion for any other function. As a further additional basis for this function only, to the extent the function solely organizes, routes, or presents information for human review, it also falls within the exclusion from the definition of automated decision-making technology for such tools (C.R.S. 6-1-1701(2)(b)(II)).

## 2.2 Generative Fallback

Where an End User prompt is not answered from Institution-Approved Content or institutional records, the AI Agent may generate a response from the underlying model's general capabilities, or indicate that it does not have a basis to answer. The substance of a Generative Fallback response is not bounded by institution-approved sources, and the model cannot reliably assess the quality or the source of its own answer. Pathify is developing a classification-and-disclosure capability (Section 5) intended to distinguish, for the End User, a response grounded in institution-approved sources from a response generated from the model's general capabilities. That capability is in development and custom rather than a current out-of-the-box feature, and its accuracy is subject to the limitation that a large language model cannot reliably self-assess.

**Classification: not Covered ADMT; an accuracy and disclosure matter.** The Generative Fallback does not make or materially influence a Consequential Decision. Its risk profile is one of accuracy and transparency rather than automated decision-making, and it is addressed through the prohibitions in Section 3, the disclosure provisions in Section 5, and, for European Union deployments, the transparency obligations of Article 50 of the EU AI Act.

## 2.3 Safety Function

Where the Customer enables it, the Safety Function uses AI-based classification to screen conversation content for indications of a possible risk of harm, such as expressions of self-harm or threats of violence. When such content is identified, the AI Agent may surface institution-designated resources, direct the End User to appropriate human contacts, and, where the Customer has configured it, route a notification to a designated Institution Designee. The classification is performed by the underlying model, which Pathify prompts to detect the relevant categories. The Safety Function is a single detection capability the institution turns on or off; it has no institution-configured categories. After classification, routing of the notification to the Institution Designee is deterministic. The notification is directed only to the Institution Designee and is not sent to the End User as an intervention; the Designee is the human in the loop between the classification and any response, receives context, and exercises independent judgment. Classifications are logged.

**Classification: a duty-of-care function, not Covered ADMT.** The Safety Function does not determine a consumer's access to, eligibility for, or the terms of an education opportunity; it is a protective and routing measure. This Policy therefore does not treat it as Covered ADMT under SB 26-189. Because Pathify offers a safety-detection capability, the function is governed instead on a duty-of-care basis: honest accuracy expectations, clear statements that it is fallible and is not a crisis or clinical service or a substitute for the institution's own safety processes, logging, and a design that supports rather than replaces the institution's response. The AI Agent does not provide mental-health treatment, crisis counseling, or emergency services, as set forth in Section 3. Consistent with Colorado House Bill 26-1263 (Conversational Artificial Intelligence Service Operator Requirements, effective January 1, 2027), the Safety Function is a detection-and-escalation and safe-messaging and referral measure: it is not crisis intervention or emergency response, it does not provide or imply licensed or certified professional care, and it is not a substitute for the institution's emergency or crisis-response systems. For European Union deployments, the Safety Function is assessed under Article 5 of the EU AI Act (emotion-inference restriction) as described in Section 10.3; Pathify's position is that text-based detection is likely outside that restriction and that the safety purpose is more likely to fit its exception than general wellbeing monitoring. As an additional basis, the Safety Function is a student-support process that does not materially influence a Consequential Decision and therefore falls within SB 26-189's exclusion from the definition of Consequential Decision for routine academic administration and student-support processes (C.R.S. 6-1-1701(3)(b)(IX)). As a further additional basis, to the extent the Safety Function solely routes or presents information for human review and the Institution Designee decides whether and how to act, it also falls within the exclusion from the definition of automated decision-making technology for such tools (C.R.S. 6-1-1701(2)(b)(II)).

## 2.4 Monitoring Function

The Monitoring Function is planned for availability. Where the Customer enables it, the Monitoring Function uses AI-based classification to identify conversations indicating institution-determined concerns, such as academic disengagement or a need for additional support, and routes a notification to a designated Institution Designee. The classification is performed by the underlying model, which Pathify prompts; unlike the Safety Function, the



Monitoring Function operates on categories the institution determines and configures; routing to the Institution Designee after classification is deterministic; the notification goes only to the Designee, who receives context and exercises independent judgment; and classifications are logged. The Monitoring Function is planned and not yet generally available; when released, it is enabled only at the Customer's election.

**Classification: a detection-and-escalation function, not Covered ADMT.** The Monitoring Function detects and escalates; it does not make or materially influence a Consequential Decision. The AI Agent classifies a conversation against institution-determined categories and routes a notification to one or more Institution Designees, who receive context and exercise independent judgment. The notification does not score, rank, prescribe, or recommend an institutional action; it surfaces the conversation for human review. The Institution Designee, not the Agent, decides whether to act on the information; nothing in the Monitoring Function auto-escalates, auto-restricts, or writes to a student record. Because the human, not the Agent, makes any consequential decision, this Policy does not treat the Monitoring Function as Covered ADMT. The Independent Determination by the Institution Designee is the central operational safeguard and the basis for this classification. If a Customer were to configure the Monitoring Function so that its output operated as a non-de-minimis factor in a Consequential Decision without a genuine Independent Determination, that configuration would fall outside this classification, and the contingent Covered ADMT procedures described in Sections 6.1, 7, 8, 8.1, 9.1, and 9.2 would apply to it. As an additional basis, the Monitoring Function addresses academic disengagement and the need for additional support and is a routine academic-administration and student-support process that does not materially influence a Consequential Decision, within SB 26-189's exclusion from the definition of Consequential Decision (C.R.S. 6-1-1701(3)(b)(IX)). As a further additional basis, because the Monitoring Function solely surfaces and routes a conversation for human review and the Institution Designee makes any determination, it also falls within the exclusion from the definition of automated decision-making technology for a tool used solely to organize, route, or present information for human review (C.R.S. 6-1-1701(2)(b)(II)).

## 2.5 Classification Summary

The following table summarizes the classification of each function. It summarizes the function classifications in this Policy, including the Integrated Action analysis detailed in Section 3.5, and is qualified by the full text of this Policy.

| Function                        | SB 26-189 treatment | Basis and governance                                                                                             |
|---------------------------------|---------------------|------------------------------------------------------------------------------------------------------------------|
| <b>Core Question-and-Answer</b> | Not Covered ADMT    | No material influence on a Consequential Decision; information-tool exclusion as secondary basis (Section 2.1).  |
| <b>Generative Fallback</b>      | Not Covered ADMT    | Accuracy and disclosure matter; prohibitions (Section 3), disclosure (Section 5), EU Article 50.                 |
| <b>Safety</b>                   | Not Covered ADMT    | Duty-of-care track (Section 2.3); honest accuracy, disclaimers, logging; EU Article 5 assessed separately.       |
| <b>Monitoring</b>               | Not Covered ADMT    | Detection and escalation (Section 2.4); an Institution Designee makes any consequential decision. The contingent |

| Function                              | SB 26-189 treatment | Basis and governance                                                                                                                                                                                                                                                                                                      |
|---------------------------------------|---------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                                       |                     | Covered ADMT procedures (Sections 6.1, 7, 8, 8.1, 9.1, and 9.2) apply only if a function is ever used to materially influence a Consequential Decision.                                                                                                                                                                   |
| <b>Integrated Action (write-back)</b> | Not Covered ADMT    | User-initiated, user-confirmed, deterministic, and Customer-configured (Section 3.5); the Agent executes the End User's confirmed instruction and does not decide the outcome. Covered ADMT analysis applies only if a configuration materially influences a Consequential Decision without an Independent Determination. |

*An Integrated Action (write-back) capability, where the Customer first configures or builds it, lets the AI Agent trigger a deterministic, user-initiated, and user-confirmed action in an institutional system, as described in Section 3.5. The Customer configures or builds the action, the End User's prompt initiates it, the End User works through and confirms a deterministic workflow in the chat, and only then is the action and its data written to a downstream system. Introducing or materially expanding such a capability requires the change-management, classification, and approval steps in Sections 9 and 12 before activation.*

## 2.6 Model Provider and Configuration

The AI Agent is powered by Google Gemini, accessed by Pathify through Google's enterprise cloud platform (the "Paid Services") under the applicable Google Cloud terms (the "Google Terms") and the Google Cloud Data Processing Addendum (the "Google DPA"), and not through the standalone Gemini API or Google AI Studio. The AI Agent reaches Gemini through the Paid Services. Pathify, not the End User, holds and uses the connection to Google.

**What is transmitted to the model provider. The AI Agent transmits to Gemini: End User prompt text (which may contain personal information the End User enters); Institution-Approved Content and other institutional context used to ground a response; public web-search results when the Customer enables that option; and, where the institution enables the use of an individual's own record context, the individual's record values themselves (for example, grades, enrollment status, or financial aid information) as needed to generate the response or perform the function. Transmissions of information directly related to an identifiable student are disclosures of FERPA Records subject to the FERPA school official architecture in Section 4.3, which for this purpose is intended to support Customer's reliance on the FERPA school official framework for transmissions to Google as the provider of the underlying model. This Policy does not represent that institution-sourced record values are abstracted away from, or withheld from, the model provider. Region-specific model-provider routing, which routes requests through regional endpoints to reduce latency and support data-residency requirements, is in development; the AI Agent currently uses the global model-provider endpoint, with regional routing expected in Q3 2026. When available, it will be managed internally by Pathify rather than configured by the institution, and it is not a customer-configurable commitment unless stated in the applicable Order Form or in data-flow documentation incorporated into the agreement.**

**Training and retention.** Pathify shall access Gemini only through the enterprise Paid Services (Google's enterprise cloud platform under the Google Terms and the Google DPA), and not through Google's consumer or no-cost services, and shall maintain those arrangements such that Google does not use Customer data or AI Agent outputs to train, develop, or improve Google's foundation models. Pathify shall maintain its configuration such that Gemini in-memory caching of AI Agent content is disabled, and shall maintain the treatment it has obtained excluding AI Agent prompts and outputs from Google's abuse-monitoring prompt logging. If Google notifies Pathify of a change to the Google Terms, the Google DPA, or the Paid Services data-handling practices that materially affects these representations, Pathify shall notify Customers within 30 days of receipt and update this Policy as appropriate. Pathify does not undertake to independently monitor Google's terms for changes; this obligation arises upon Pathify's receipt of Google's notice or actual discovery of a material change.

**Minors and the model provider.** The AI Agent is not designed for or directed to minors: Pathify, not the student, holds and uses the model-provider connection, and the system is not targeted to minors. Incidental under-18 use (for example, dual-enrollment, pre-matriculation, or pre-college users) is foreseeable. Customers remain solely responsible, as between the parties and consistent with the Google Terms, for child-privacy, online-safety, and age-appropriate-design compliance applicable to their deployments.

## **2.7 Product Surface and Activation Boundaries**

This Policy applies to the AI Features Pathify expressly makes available for Customer production use under the governing Order Form and Customer configuration. Beta, preview, pilot, experimental, separately ordered, or materially different AI-enabled features are not authorized for Customer production use unless expressly enabled by Pathify and the Customer through the applicable contractual and configuration process. Where Pathify introduces a materially different AI feature, model surface, orchestration component, or data-flow pattern, that feature will be evaluated under this Policy before production activation, and Pathify will determine whether corresponding updates to this Policy, the DPA, the Terms of Use, or the CXP Privacy Policy are required.

## **3. Prohibited Uses**

Neither Pathify, any Customer, nor any End User shall use or configure the AI Agent to do any of the following in contravention of this Policy.

### **3.1 Consequential Decision Uses (Core Question-and-Answer Function)**

Generated content of the AI Agent's Core Question-and-Answer Function shall not be used to make or materially influence a Consequential Decision (as defined in C.R.S. 6-1-1701(3)), including any Consequential Decision within the education covered domain under C.R.S. 6-1-1701(6)(a), such as a decision concerning education enrollment or an education opportunity. This prohibition keeps the Core Question-and-Answer Function from materially influencing a Consequential Decision and, as a secondary and additional basis for that function only, supports the ADMT Information-Tool Exclusion under C.R.S. 6-1-1701(2)(b)(III). It does not limit the separate treatment of the Safety and Monitoring detection-and-escalation functions under this

Policy. Without limiting the foregoing, the Core Question-and-Answer Function shall not be configured or used to:

- Advise students on course selection, degree requirements, or graduation pathways in a manner that constitutes or materially influences an institutional determination.
- Generate a student-specific recommendation that functions as a factor in admissions, financial aid, academic standing, placement, registration, or graduation.
- Produce an output represented to the End User as an institutional decision rather than as information or a referral.

### **3.2 Discriminatory or Harmful Content**

The AI Agent shall not be used to generate content that unlawfully discriminates against any individual or group on the basis of a protected characteristic under applicable federal, state, or local anti-discrimination law, or to generate harassing, threatening, or otherwise unlawful content. This prohibition applies to every function.

### **3.3 Professional-Advice Uses**

The AI Agent shall not be used to provide medical, legal, financial, clinical, or other professional advice, diagnosis, or treatment. It may refer End Users to institutional resources and qualified professionals.

### **3.4 Excluded AI Uses**

The AI Agent shall not be used, and the Customer shall not configure or permit End Users to use the AI Agent, to provide or simulate mental health services, crisis intervention, or emergency response (collectively, “Excluded AI Uses” as defined in the MSA). These are high-stakes human-service functions requiring qualified professional intervention, direct human judgment, and professional liability that the AI Agent is not designed, trained, or authorized to provide. The Safety Function (Section 2.3) routes a notification to a human Institution Designee and is not itself an Excluded AI Use, but it is not a substitute for the institution’s own crisis and emergency processes. Customer is solely responsible for ensuring End Users have access to appropriate professional resources for Excluded AI Uses through channels other than the AI Features.

### **3.5 Integrated Action (Write-Back)**

The AI Agent can trigger a deterministic action in the Customer’s or a connected institutional system from within the AI Agent user experience (an “Integrated Action,” also called write-back). A Customer configures or builds the Integrated Actions available to its End Users, including by building deterministic write-backs. Customer-built deterministic write-backs are available today, subject to the four conditions below; additional pre-built actions are offered as they are released. Examples of Integrated Actions include submitting an assignment or scheduling an advising appointment. An Integrated Action is available only where, and to the extent, the Customer has first configured or built it; no Integrated Action is available by default.

Pathify designs Integrated Actions to meet four conditions: (i) it is user-initiated, meaning the End User, not the Agent on its own initiative, starts the action; (ii) it is user-confirmed, meaning

the AI Agent presents a distinct confirmation step that shows the specific action and its consequence, and the End User affirmatively confirms that action before it executes; an initiating click or request alone is not sufficient confirmation; (iii) the executed write-back is deterministic, meaning that when an action runs, the Agent triggers a predefined integration and does not itself decide the substantive outcome or generate the action freely; the conversation that leads to the action may be generative, but the write-back it triggers is not; and (iv) it is Customer-configured, meaning only Integrated Actions the institution has first configured or built are available, and none is available by default. Because an Integrated Action carries out the End User's own confirmed instruction through a deterministic integration, the Agent executes the user's request rather than making a Consequential Decision, and this Policy does not treat the availability of Integrated Actions, by itself, as Covered ADMT. Pathify's pre-built Integrated Actions meet these conditions by design. A Customer may also build its own write-backs, and Pathify does not architecturally force these conditions on Customer-built write-backs. The Customer, as deployer, is responsible for ensuring that any write-back it builds is user-initiated, presented to the End User for confirmation of the specific action and its consequence, and deterministic, and for not building a write-back through which the Agent's output materially influences a Consequential Decision without a genuine Independent Determination. A Customer that builds a write-back operating outside these conditions is acting as the deployer of that action, and the analysis in Section 1.2, the contingent Covered ADMT procedures, and the Customer's own obligations under applicable law apply to that configuration.

Where an Integrated Action transmits or writes personal data, that data flow is subject to Section 4 and, for education records, the FERPA school official architecture in Section 4.3. Introducing a new Integrated Action, or materially expanding the categories of data written or the systems affected, is a material change requiring the change-management, classification, and approval steps in Sections 9 and 12 before activation.

## **4. Data Handling**

### **4.1 Data Transmitted to the Model Provider**

The following table describes what is and is not transmitted to Google Gemini. Whether an individual's own record values are transmitted depends on the institution's configuration of the Core Question-and-Answer Function and on activation of the Monitoring Function. As of the date of this Policy, the AI Agent's production data flow to the model provider is limited to the categories marked as transmitted in this Section; no record value reaches the model provider in production except where the institution has configured the Core Question-and-Answer Function to use an individual's own record context or has activated the Monitoring Function. When the User Context (Section 4.5), retrieval of an individual's own records from institutional systems, or other record-context features are released and enabled, configured record values may be transmitted as described in this table.

| Data category                                       | Transmitted to Gemini? | Notes                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
|-----------------------------------------------------|------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Conversation text (prompts)                         | Yes                    | End User's typed questions and prompts, which may include personal information the End User enters.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| Institutional content                               | Yes                    | Published pages, FAQs, directories, and other Institution-Approved Content used to ground a response. Where the Customer enables connected sources such as SharePoint, content retrieved from those sources is grounded the same way.                                                                                                                                                                                                                                                                                                                                                              |
| Public web-search results                           | Yes (when configured)  | Only when the Customer enables web search.                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
| Record values / personally identifiable information | Yes (when configured)  | When the institution enables the use of an individual's own record context (for example, grades, enrollment status, financial aid), the corresponding record values are made available to the model provider only as needed to generate the response or perform the institution-enabled function, subject to the FERPA school official architecture in Section 4.3. The User Context feature is on by default; however, each data point is off by default, and no record value reaches the model provider through the User Context unless and until the Customer enables that specific data point. |
| Voice input                                         | As transcribed text    | Where voice is supported, audio is transcribed to text and only the text is processed; the audio signal itself is not analyzed (Section 4.6). A live-voice capability that processes the voice signal itself is a distinct, planned capability addressed separately in Section 4.6 and is not part of the current transcription-only handling described here.                                                                                                                                                                                                                                      |
| User-uploaded files                                 | Yes, when enabled      | Files an End User uploads in a conversation, where the Upload File feature is enabled; transmitted to the model provider as part of the input. Planned.                                                                                                                                                                                                                                                                                                                                                                                                                                            |

Key: cells shaded green indicate data categories transmitted to the model provider in ordinary operation; cells shaded amber indicate categories transmitted only when the institution enables that specific category.

## 4.2 Record Retention

- **Conversation transcripts.** Retained for the duration of the Customer contract. Accessible by institutional administrators and by authorized Pathify staff with a need to know, subject to role-based access controls and logging.
- **Internal logs.** Retained for thirty days as a rule; may contain prompts in error messages; accessible only by Pathify staff.



- **Activity records (in development).** Separately from the diagnostic logs described above, Pathify records actions taken in the platform, including changes to configuration settings and End User actions such as page visits and interactions, without the content of those interactions. These records support the analytics available to institutional administrators. Pathify is implementing a retention schedule for these records under which a record of a configuration change, including the acting administrator and the time of the change, is retained for up to ten years from the date it is recorded, and all other activity records are deleted within ninety days following termination or expiration of the Customer contract. Until that schedule is implemented, these records are retained without a defined expiry; the schedule will apply to records already held when it takes effect.
- **Classification logs.** Safety and Monitoring classifications are logged with the input and the outcome, supporting detection-and-escalation defensibility and the recordkeeping that would support a Covered ADMT classification if one were ever to apply.
- **FERPA considerations.** To the extent transcripts or logs contain information directly related to an identifiable student and are maintained by or on behalf of an institution, they may constitute FERPA Records. Customers are responsible for determining applicable retention periods under their own FERPA policies and record-retention schedules.

#### 4.3 FERPA School Official Architecture

Where the AI Agent transmits information directly related to an identifiable student, whether to the model provider to generate a response or perform a function, or to an Institution Designee through a referral, Safety notification, or Monitoring notification, those transmissions may constitute disclosures of FERPA Records and are intended to be made in reliance on the school official exception under 34 C.F.R. 99.31(a)(1)(i). Pathify and the Customer accordingly agree as follows. For clarity, information directly related to an identifiable student may reach the model provider through conversation text entered by an End User, independently of whether the Customer has enabled the User Context under Section 4.5.

- **Customer designations.** The Customer shall designate each Institution Designee as a school official under the Customer's FERPA policy and shall determine, for each Designee category, the scope of legitimate educational interest applicable to AI-Agent-originated referrals, classifications, and evaluations. The Customer is responsible for the reasonable methods required by 34 C.F.R. 99.31(a)(1)(ii). The per-data-point controls described in Section 4.5 are available to the Customer for that purpose.
- **Annual notification.** The Customer represents and covenants that, to the extent any AI Agent workflow involves transmission of information directly related to an identifiable student, its FERPA annual notification under 34 C.F.R. 99.7 reflects, or will reflect before activation of the applicable function, the school official designation as including AI-Agent-originated transmissions, and that the notification's description of criteria for legitimate educational interest is consistent with the scope determined under this Section 4.3.
- **Pathify FERPA status.** Pathify's handling of FERPA Records in connection with the AI Agent is subject to the school official and contractor terms in the MSA (including the terms addressing transmissions to the provider of the underlying model) and the DPA, including

direct control, use and redisclosure restrictions, and the requirements of 34 C.F.R. 99.31(a)(1)(i)(B).

- **Google as model provider.** Where the AI Agent transmits FERPA Records to Google as the provider of the underlying model, Pathify shall access Gemini only through Google's enterprise Paid Services under the Google Terms and the Google DPA referenced in Section 2.6, and not through Google's consumer or no-cost services, and shall maintain contractual arrangements with Google to support Customer's reliance on the FERPA school official framework for those transmissions, relying on the processor-controller framework and use and confidentiality restrictions of the Google DPA as applied to FERPA Records. Pathify shall maintain those arrangements such that, as applied to FERPA Records, Google processes them only to provide the Paid Services to Pathify and the Customer, under the use, confidentiality, and redisclosure restrictions and the direct-control requirements of 34 C.F.R. 99.31(a)(1)(i)(B), and not for Google's own purposes. Pathify will provide Customer, on request, summary documentation of the relevant Google provisions. If Google amends the Google Terms or Google DPA in a manner Pathify reasonably determines materially affects the FERPA framework, Pathify shall notify the Customer within 30 days after the amendment takes effect and shall work with the Customer to identify any appropriate revisions to this Policy.
- **Disclosure log.** The Customer is responsible for maintaining any disclosure record required under 34 C.F.R. 99.32 to the extent AI-Agent-originated transmissions trigger that obligation. Pathify shall provide reasonable cooperation and the transmission metadata reasonably necessary to support it.

***Important FERPA limit. De-identification under FERPA is stricter than name removal, and cross-institutional pooling of student data without express authorization from each institution is not permitted. The User Context (Section 4.5), and any data made available to the model, are scoped to the individual Customer institution; the AI Agent is not to create cross-institutional data flows.***

#### **4.4 Safety and Monitoring Classification (Operational Detail)**

This Section 4.4 describes the operational detail of the Safety Function (Section 2.3) and the Monitoring Function (Section 2.4). For both, classification is performed by the underlying model, which Pathify prompts to detect the relevant categories. The Safety Function has no institution-configured categories; it is turned on or off. The Monitoring Function operates on categories the institution determines and configures. For each, routing of a notification to the designated Institution Designee after classification is deterministic; the notification is directed only to the Institution Designee and not to the End User as an intervention; the Institution Designee receives context, no recommended action, and exercises independent judgment; and classifications are logged. The institution's designated staff are responsible for reviewing flagged conversations and determining the appropriate response. The institution remains the controller of the information consistent with the rest of this Policy and FERPA.

**Customer option and End-User notice.** The institution chooses whether to enable Safety detection and, when available, Monitoring detection. Where enabled, the AI Agent uses automated detection, which is AI-based and fallible and is not literal keyword matching: it may

flag conversations that do not contain specific trigger words and may fail to flag conversations that do. It is not a replacement for the institution's safety processes and protocols. If an End User is in crisis or needs immediate help, the End User should contact the institution's emergency or crisis-response service or the emergency services in the relevant jurisdiction. Where the AI Agent identifies a relevant conversation, the conversation is recorded in Conversation History accessible to the institution's designated administrators, and a notification may be sent to specific institutional staff the institution designates.

The Safety and Monitoring functions are designed to use the same data-flow envelope otherwise applicable to the AI Agent's active configuration and not, by themselves, to expand the categories of personal data transmitted to the model provider beyond what Section 4.1 describes.

#### **4.5 Institution-Configured User Context**

The institution-configured user context (the "User Context") is an institution-scoped capability that allows the AI Agent to use an individual's own records held in the Customer's systems, whether that individual is a student, a staff member, a prospective student, or another individual whose records the Customer has made available, within the Customer's environment to respond in a more personalized way, including by mapping an inquiry to the appropriate institutional resource, help-desk function, or advisor. The User Context governs the record values themselves, without regard to which End User is interacting with the AI Agent in a given session; where the Customer permits an End User to access another individual's records, this Section governs whether those record values are made available to the AI Agent and to the model provider. The User Context feature is on by default; however, each data point is off by default, and no record value is included in the context made available to the AI Agent, or transmitted to the model provider, unless and until the Customer enables that specific data point. The Customer determines which record values are included on a data-point-by-data-point basis, with each data point (for example, grade point average and student identifier) separately toggleable on its own, and each configuration change is logged. To the extent the institution has enabled specific data points, the corresponding record values are made available to the AI Agent only as needed to perform the institution-enabled function or generate the requested response. The User Context, and any data it makes available, is scoped to the individual Customer institution; it is an architectural requirement that one Customer's data not be accessible to, or used for the benefit of, another Customer, and that the AI Agent not create cross-institutional data flows. Where the institution has enabled data points that include a student's record values for use in generating a response, those values are made available to the model provider as described in Sections 2.6 and 4.1 and are subject to the FERPA school official architecture in Section 4.3. Record values of individuals who are not students are not FERPA Records, and are handled in accordance with the Data Processing Addendum and the CXP Privacy Policy. A change that would materially expand the personally identifiable context made available to the model provider is a material change in data flow requiring Customer notice under Section 12.

#### **4.6 Voice Interactions**

Dictation (speech-to-text) is scheduled to enter production at the end of July 2026. Live voice interaction (real-time voice signal analysis) is a distinct, later capability planned for a future release. Where the AI Agent supports voice interaction, audio input is transcribed to text and only the resulting text is processed; the AI Agent operates on the transcript, not on the voice signal itself, and Pathify does not analyze vocal characteristics. Pathify does not intentionally retain audio recordings in the AI Agent workflow, and the transcript is handled in the same manner as text-entered prompts under this Section 4. Because the AI Agent works only on transcribed text, voice input is in the same position as typed text for the European Union emotion-inference analysis in Section 10.3; this position holds only while voice remains transcription-only. Pathify is planning a live-voice capability in which the AI Agent would process the voice signal itself rather than only a transcript. That capability is distinct from the transcription-only handling described above, would process biometric data, and would require a separate analysis under the European Union emotion-inference restriction (Section 10.3) and applicable biometric-data and consent requirements. Where it is offered, Pathify may make it available only in configurations or jurisdictions in which it can be offered consistent with Applicable Laws, and it may be excluded from European Union and United Kingdom deployments. Until that capability is enabled for a given configuration, voice remains transcription-only as described above. The Customer remains responsible for any voice-recording notice or consent required under applicable law for its configuration and use.

#### **4.7 Data-Flow Documentation**

Pathify will maintain, and make available to Customers on request, reasonable documentation describing each AI Feature's material data-flow characteristics, including the applicable model-provider surface, the categories of Customer-designated data sources used, whether data is used for grounding, retrieval, support, monitoring, or security functions, the retention posture Pathify understands to apply, and the customer-facing controls available. Pathify may satisfy this through this Policy, the DPA, a Product Surface and Data-Flow Schedule, an AI Feature Authorization Schedule, or other written governance documentation. Nothing in this Section freezes Pathify's architecture or prevents future updates made in accordance with Section 12.

#### **4.8 Operational Records and Evidence**

Where supported by the applicable AI Feature and configuration, Pathify may maintain logs, metadata, transcript records, activation records, configuration records, testing records, or other operational evidence relating to AI Feature use for security, abuse prevention, troubleshooting, governance, compliance support, and product operation. Retention, granularity, and exportability may vary by feature and configuration. If a function were ever classified as Covered ADMT, Pathify and the Customer would maintain the records described in Sections 6, 7, and 8 for it.

#### **4.9 Accessibility**

Pathify designs the AI Agent's End User interface to be usable by people with disabilities and to support Customers' obligations under the Americans with Disabilities Act (ADA), Section 504 of the Rehabilitation Act, and, for public institutions, Title II of the ADA and the related United States Department of Justice web and mobile accessibility regulation. Pathify's accessibility

target for the AI Agent interface is conformance with the Web Content Accessibility Guidelines (WCAG) 2.2 Level AA.

On request, Pathify will provide available accessibility conformance documentation describing the AI Agent's conformance and any known limitations, including a current accessibility conformance report (for example, a Voluntary Product Accessibility Template, or VPAT) where one is available, and will work with the Customer on a reasonable remediation timeline for material accessibility defects identified in the AI Agent interface. Conformance is assessed against these commitments rather than against a requirement that every interface element conform at every moment.

## 5. AI Disclosure

SB 26-189 imposes a point-of-interaction notice under C.R.S. 6-1-1704(1) on a deployer using Covered ADMT to materially influence a Consequential Decision; it does not impose a standalone requirement that consumers be told they are interacting with an AI system for non-covered functions. Independent of that trigger, this Policy requires a clear disclosure that the End User is interacting with an AI system and not a human, on the terms set out in this Section 5 and Section 5.1. Pathify treats this as a baseline practice that supports the AI Agent's informational positioning and, for the Core Question-and-Answer Function, the ADMT Information-Tool Exclusion, and that aligns with Article 50 of the EU AI Act for European Union deployments.

Opening-message disclosure. Pathify shall ensure the AI Agent's user interface includes an explicit disclosure that the End User is interacting with an AI system, appearing in the opening message at the start of each interaction.

**Answer-source marker (in development).** Pathify is developing a marker that distinguishes, for the End User, a response grounded in Institution-Approved Content or institutional records from a Generative Fallback response generated from the model's general capabilities. This marker is in development and custom rather than a current out-of-the-box feature, and a large language model cannot reliably self-assess the quality or the source of its own answer, so the marker's reliability is inherently limited. Pathify will not represent the marker as more reliable than it is, and will describe its limitations where it is offered.

**Persistent interaction disclosure. Where the Customer has activated the AI Features, the AI Agent displays a persistent disclosure with the prompt-entry field, in substantially the following form:**

*"This is an AI assistant powered by Google Gemini. It may respond with inaccurate or harmful information."*

Customers shall not remove, obscure, or disable this disclosure without providing a substitute that is at least equally clear and that satisfies any applicable AI disclosure obligations, including, where the Customer uses Covered ADMT to materially influence a Consequential Decision, the point-of-interaction notice required by C.R.S. 6-1-1704(1).

### 5.1 Acceptance Screen and Model Disclosure

Acceptance screen. Before an End User may continue to use the platform, the Customer's deployment presents a one-time acceptance screen requiring the End User to acknowledge having read this Policy, the CXP privacy policy published by the Customer, and the Terms of Use. The acceptance screen includes an AI disclosure block, drafted by Pathify, stating that the End User is interacting with an artificial intelligence system and not a person, identifying the model provider, and stating that institutional staff, not the AI Agent, make decisions about admission, enrollment, financial aid, grades, discipline, and similar outcomes. Where the Customer uses Covered ADMT to materially influence a Consequential Decision, the acceptance screen also carries the point-of-interaction notice required by C.R.S. 6-1-1704(1).

Pathify records the fact, version, and timestamp of each End User's acknowledgment. The Customer may add institution-specific content to the acceptance screen, including the Customer's own usage terms, but shall not remove the AI-interaction disclosure, the model-provider disclosure, or the statement that institutional decisions are made by institutional staff, and may reword them only where the substance of each is preserved. Each End User will be required to acknowledge the acceptance screen, or to acknowledge it again, before continuing to use the platform.

## 6. AI Impact Analysis

Pathify maintains an AI impact analysis for the AI Agent and provides it to Customers as part of their AI-governance and procurement review, regardless of how any single function is classified. This Section assesses the AI Agent against each category of Consequential Decision within the covered domains enumerated in SB 26-189 (C.R.S. 6-1-1701(6)). The analysis distinguishes the Core Question-and-Answer Function, which does not materially influence a Consequential Decision, from the Monitoring Function, which is a detection-and-escalation function where an Institution Designee, not the Agent, makes any consequential decision. As of the date of this Policy, no AI Agent function is Covered ADMT, and Pathify maintains this impact analysis as a standing governance artifact it provides to Customers, with its formal status as a Covered ADMT impact assessment attaching only if a function were ever classified as Covered ADMT. SB 26-189 does not itself mandate an impact assessment; Pathify provides this framework to support Customers' own AI governance and any impact-assessment processes they undertake.

| Consequential Decision category     | Core Q&A impact | Basis                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    |
|-------------------------------------|-----------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Education enrollment or opportunity | <b>None</b>     | Core Q&A surfaces published information and the individual's own records and routes inquiries; it does not make or influence enrollment, admissions, aid, placement, registration, or graduation decisions. Section 3.1 prohibits such configuration. Monitoring is a detection-and-escalation function where an Institution Designee makes any consequential decision; it is not Covered ADMT. The contingent procedures in Sections 6.1, 7, and 8 would apply only if a function were ever classified as Covered ADMT. |



| Consequential Decision category      | Core Q&A impact | Basis                                                                                                                                     |
|--------------------------------------|-----------------|-------------------------------------------------------------------------------------------------------------------------------------------|
| Employment or employment opportunity | None            | Not used in hiring, evaluation, termination, or any employment decision.                                                                  |
| Financial or lending service         | None            | Does not evaluate creditworthiness or determine loan eligibility; may surface published financial-aid information.                        |
| Essential government service         | None            | Deployed in higher education; no interaction with government benefit or eligibility determinations.                                       |
| Healthcare services                  | None            | Prohibited from providing medical advice, diagnosis, treatment, or counseling (Section 3.3); may refer to institutional health resources. |
| Housing, insurance, legal service    | None            | No role in housing assignment or eligibility, insurance, or the provision of legal advice (Sections 3.3).                                 |

## 6.1 Impact Documentation for Covered Functions

No AI Agent function is currently Covered ADMT, so this Section requires no impact documentation today. If a function were ever classified as Covered ADMT, Pathify and the Customer would maintain impact documentation for it identifying: (i) the categories the function detects or evaluates; (ii) the inputs the function processes; (iii) the Institution Designee category and the scope of Independent Determination authority; (iv) the notifications delivered and to whom; and (v) the anticipated effect, if any, on the Consequential Decision categories above. For an education Consequential Decision, a Customer subject to FERPA may rely on its existing FERPA notices and student-record procedures to satisfy the related notice and consumer-rights obligations (Sections 8 and 8.1).

## 6.2 Impact Support for Public Higher Education Procurement

This framework also supports institution-side procurement and governance review, including in jurisdictions such as New York that expect vendors to help institutions document objectives, system boundaries, data categories, foreseeable misuse, privacy and cybersecurity dependencies, user-notification implications, and human-review points. For each Covered ADMT function, the documentation should identify: the objective; the AI workflow at a functional level; the categories and sources of data; foreseeable misuse or failure modes; privacy, FERPA, and cybersecurity dependencies; user-notification implications; the human-review or determination point; and the threshold for re-assessment after a material change.

## 7. Algorithmic Discrimination Risk Assessment

Pathify treats non-discrimination and equitable treatment of users as a standing commitment for the AI Agent and maintains the assessment in this Section regardless of statutory classification. SB 26-189 does not impose a standalone duty to prevent “algorithmic discrimination”; the former Colorado Artificial Intelligence Act’s algorithmic-discrimination provisions were repealed, and SB 26-189 instead allocates, between developers and deployers, fault for unlawful discrimination under existing law. For this Policy, “algorithmic discrimination” means a condition in which the

use of an AI system results in unlawful differential treatment or impact disfavoring an individual or group on the basis of a protected characteristic under applicable anti-discrimination law. Pathify maintains this discrimination-risk assessment as a standing safeguard. For the foundation model, Pathify relies on the model provider's published model documentation and evaluations (for example, Google's model cards) rather than conducting its own foundation-model bias testing; the Agent's informational, non-consequential design is the primary reason its bias exposure is low. Pathify provides this assessment to Customers to support their own bias-evaluation obligations.

Risk profile by function:

- **Core Question-and-Answer and Generative Fallback.** Outputs are informational rather than decisional, which limits the pathway through which a discriminatory output could produce a material legal or similarly significant effect. Section 3.2 prohibits discriminatory content. Pathify monitors for patterns of discriminatory output as part of ongoing product-quality processes.
- **Monitoring Function (detection and escalation, not Covered ADMT).** Because the Monitoring Function is not currently Covered ADMT, the Covered ADMT-specific discrimination-risk assessment obligations do not apply to it today; Pathify nonetheless applies the standing non-discrimination and product-quality safeguards described in this Section. If the Monitoring Function or any other function were ever classified as Covered ADMT, Pathify and the Customer would assess whether the detection categories and inputs create a known or reasonably foreseeable risk of differential treatment or impact on the basis of a protected characteristic, and would identify mitigations, including, where feasible given the Customer's data, pre-activation testing of the classification for disparate impact across relevant demographic groups; ongoing review of Institution Designee escalation patterns for indications of embedded bias; and review of detection criteria for inputs that may proxy for protected characteristics. The assessment would be documented and made available to the Customer on request.

For a function classified as Covered ADMT, or as the Monitoring Function is finalized for general availability, the assessment for that function identifies the testing scope, sample inputs or benchmark scenarios where feasible, the review cadence, the personnel responsible, the escalation threshold, and the remediation steps, satisfied through an internal protocol, a third-party review, or a hybrid that is documented and consistent with the sensitivity of the use case. This function-level assessment is distinct from foundation-model bias, which Pathify addresses through the model provider's published evaluations as described above. Pathify does not represent that it conducts application-level bias testing for a function that is not currently Covered ADMT; the depth and independence of any such testing for the Monitoring Function is a matter Pathify will confirm when finalizing that function for general availability.

## 8. Developer and Deployer Role Allocation

Under SB 26-189, obligations differ based on whether an entity is a "developer" (C.R.S. 6-1-1701(8)) or a "deployer" (C.R.S. 6-1-1701(7)) of Covered ADMT. These obligations attach

with respect to Covered ADMT, that is, automated decision-making technology used to materially influence a Consequential Decision.

- **Pathify: developer.** Pathify develops and maintains the AI Agent. No AI Agent function is currently Covered ADMT, so Pathify bears no developer obligations under C.R.S. 6-1-1702 today. If a function were ever classified as Covered ADMT, Pathify would bear developer obligations under C.R.S. 6-1-1702, principally making available to each deployer documentation describing the Covered ADMT's intended uses, the categories of data used to train it (to the extent known), its known limitations, and instructions for appropriate use and meaningful human review, together with notice of material updates and three-year record retention. Under C.R.S. 6-1-1707, a developer is liable for unlawful discrimination involving Covered ADMT only to the extent the deployer used the Covered ADMT in a manner intended, documented, marketed, advertised, configured, or contracted for by the developer; use of the AI Agent outside its documented and configured uses is not a developer-authorized use.
- **Customer: deployer.** The Customer deploys the AI Agent to its community through its configured instance of the Pathify platform. No AI Agent function is currently Covered ADMT, so the Customer bears no deployer obligations under SB 26-189 for the AI Agent today. If a function were ever classified as Covered ADMT, the Customer would bear deployer obligations under SB 26-189, principally the point-of-interaction notice (C.R.S. 6-1-1704(1)), the post-adverse-outcome disclosure within 30 days (C.R.S. 6-1-1704(3)), the consumer rights of correction and of meaningful human review and reconsideration (C.R.S. 6-1-1705), and three-year record retention (C.R.S. 6-1-1703). For an education Consequential Decision, a Customer subject to FERPA may satisfy these notice, disclosure, correction, and human-review obligations through its existing FERPA notices and student-record inspection, review, and amendment procedures (C.R.S. 6-1-1704(9) and 6-1-1705(2)), without a separate or duplicative process. As deployer, the Customer is also responsible for the meaningful human oversight of the conversations and notifications the AI Agent surfaces through the Safety and Monitoring Functions, including ensuring that the institutional staff who review them exercise independent judgment rather than deferring automatically to the Agent, and for any Integrated Action it builds as described in Section 3.5.
- **Core Question-and-Answer Function.** Because the Core Question-and-Answer Function is not Covered ADMT, the developer and deployer obligation frameworks are not triggered for it. As a secondary basis for that function only, the configuration also supports the ADMT Information-Tool Exclusion (Sections 2.1 and 3.1).

The contractual allocation of developer and deployer obligations between Pathify and the Customer is set forth in the Colorado and Similar AI Law Allocation section of the MSA.

### 8.1 Consumer Rights and Adverse-Outcome Procedures (Covered Functions)

This Section applies to any AI Agent function that is classified as Covered ADMT and to any configuration in which the AI Agent materially influences a Consequential Decision. No AI Agent function is currently Covered ADMT, so this Section imposes no current obligation. It does not apply to the Core Question-and-Answer Function. These are obligations of the Customer as

deployer; Pathify supports the Customer but does not assume them and does not guarantee the Customer's compliance.

- **Point-of-interaction notice.** Before using a Covered ADMT function to materially influence a Consequential Decision, the Customer shall provide the affected consumer a clear and conspicuous notice consistent with C.R.S. 6-1-1704(1) and (2). Pathify will make available functionality and template language reasonably necessary to support this notice.
- **Post-adverse-outcome disclosure.** If a Covered ADMT function materially influences a Consequential Decision resulting in an adverse outcome, the Customer shall provide, within 30 days, (i) a plain-language description of the decision and the role the Covered ADMT played, (ii) instructions and a simple process to request additional information (including the name and version of the AI Agent, its developer, and the types, categories, and sources of personal data used), and (iii) an explanation of the consumer's rights, consistent with C.R.S. 6-1-1704(3). Pathify will make available the developer-side information in Section 8 reasonably necessary to populate item (ii).
- **Correction, human review, and reconsideration.** On request following such an adverse outcome, the Customer shall provide instructions for correcting factually incorrect or materially inaccurate personal data used in the decision, and an opportunity for meaningful human review and reconsideration to the extent commercially reasonable, consistent with C.R.S. 6-1-1705. Correction is not required for opinions, predictions, scores, or protected evaluations.
- **FERPA process alignment for education decisions.** For an education Consequential Decision, a Customer subject to FERPA may satisfy the notice, disclosure, correction, and human-review obligations of this Section through its existing FERPA notices and student-record procedures, including any applicable complaint or appeal process, without a separate or duplicative process, consistent with C.R.S. 6-1-1704(9) and 6-1-1705(2). This is what would make the contingent Covered ADMT consumer-rights obligations manageable in practice for FERPA-subject institutions if a function were ever so classified. Consistent with SB 26-189, a FERPA-covered educational institution that follows its FERPA notice and disclosure requirements is treated as compliant with the corresponding obligations of SB 26-189.
- **Adverse-outcome runbook.** Before activating any function as Covered ADMT, Pathify will maintain an internal runbook for supporting Customers when a Covered ADMT function is implicated in an adverse outcome, including the developer information to be supplied and the timing, and a consumer-rights assist mechanism through which a Customer can request that information. These are operational build items Pathify is implementing as the Monitoring Function moves to general availability.

**Relationship to the DPA.** Where a Customer has activated a Covered ADMT function, the procedures in this Section apply and the DPA incorporates this Section by reference to the extent the procedures involve processing of personal data. Where the Customer has activated no Covered ADMT function, this Section imposes no operative obligation.

## 9. Function Status, Activation, and Change Management

Function status is determined by what is built and enabled, not by a calendar date or a tier label.

- **Core Question-and-Answer Function.** Generally available and active. Covers Institution-Approved Content retrieval, answering questions, navigation, display of an individual's own records where configured, and informational referral.
- **Safety Function.** Available; enabled only at the Customer's election. Governed on the duty-of-care basis in Section 2.3.
- **Monitoring Function.** Planned for availability; enabled only at the Customer's election; a detection-and-escalation function, not Covered ADMT, in which an Institution Designee makes any consequential decision. The contingent procedures in Sections 6, 7, and 8 would apply only if the function were ever classified as Covered ADMT.
- **Integrated Action (write-back) within institutional systems.** A Customer-configured capability, available only where the Customer first configures or builds it (for example, submitting an assignment or scheduling an advising appointment), and subject to the four conditions in Section 3.5.

## 9.1 Activation Conditions for Covered Functions

No current AI Agent function is Covered ADMT, so the following activation conditions impose no obligation today. If a function were ever classified as Covered ADMT, the following conditions would be completed before that function is used to materially influence a Consequential Decision:

1. Pathify, as developer, provides the developer documentation under C.R.S. 6-1-1702 for the function.
2. Pathify and the Customer complete the impact documentation (Section 6.1) and the algorithmic discrimination risk assessment (Section 7) for the function.
3. The Customer designates Institution Designees and the scope of Independent Determination authority, and confirms the FERPA school official designation extends to the AI-Agent-originated transmissions (Section 4.3).
4. The Customer is notified of its deployer obligations under SB 26-189 (point-of-interaction notice, post-adverse-outcome disclosure, and correction and meaningful-human-review rights), which a FERPA-subject Customer may satisfy through existing FERPA processes (Section 8.1).
5. The MSA's Colorado and Similar AI Law Allocation section is updated, or a Covered ADMT Approval Instrument (as defined in the MSA, meaning a supplemental Order Form or written MSA update) is executed, to allocate the developer and deployer obligations for the function.

## 9.2 Risk Management for Covered Functions

As a standing governance practice, Pathify maintains product-level risk-management practices for the AI Agent that are structured to align with the NIST AI Risk Management Framework (NIST AI RMF) and ISO/IEC 42001; since Pathify supports each Customer's own risk-management program, this reflects alignment and an ongoing diligence commitment rather than a representation of compliance with any specific framework. The deployer-side controls

described below become contractually required only if a function is ever classified as Covered ADMT. SB 26-189 does not require a risk management program. As a condition of activating a Covered ADMT function, however, Pathify requires the Customer, as deployer, to maintain a risk management policy and program for that function that specifies the principles, processes, and personnel used to identify, document, and mitigate known or reasonably foreseeable risks of unlawful discrimination; is iterative over the life cycle; and is reasonable considering the NIST AI RMF, ISO/IEC 42001, or an equivalent framework, the size and complexity of the deployer, the nature and scope of the function, and the sensitivity and volume of data processed. The policy should address a written governance statement tied to the NIST AI RMF core functions (Govern, Map, Measure, Manage) or ISO/IEC 42001 controls; designated accountable personnel; a documented process for evaluating discrimination risk including pre-deployment bias testing and post-deployment outcome review; incident-response procedures; impact documentation retained with the deployer records under C.R.S. 6-1-1703 for not less than three years; and consumer notice and rights procedures implementing C.R.S. 6-1-1704 and 6-1-1705, which a FERPA-subject Customer may satisfy through existing FERPA processes.

Upon activation of a Covered ADMT function, Pathify makes available the developer documentation required under C.R.S. 6-1-1702(1), provides notice of material updates under C.R.S. 6-1-1702(2), and retains related records for not less than three years under C.R.S. 6-1-1702(4).

### **9.3 Change Management and Architectural Principle**

The organizing principle is that the AI Agent informs, classifies, refers, and, where the institution configures or builds an Integrated Action, triggers a deterministic action the End User has initiated and confirmed, while a qualified human makes every Consequential Decision. The AI Agent should never itself decide an outcome that affects a student's enrollment, financial aid, academic standing, or similar status; where an Integrated Action carries out such a step, it does so only on the End User's own confirmed instruction. Before deploying any new feature that would materially expand the AI Agent's functions, Pathify shall: (i) identify the function and its classification under this Policy; (ii) confirm that system architecture, not policy language alone, enforces the boundary that the AI Agent's output cannot trigger an institutional action without either a genuine Independent Determination where a Consequential Decision is implicated or a distinct End User confirmation where the action is a user-initiated Integrated Action; and (iii) determine which companion instruments require update before the feature ships. A change in whether record values or other personal data are transmitted to the model provider requires updates to Sections 2.6 and 4.1, the DPA, and the FERPA analysis. A change that alters a function's regulatory classification requires updates to Sections 2, 6, and 8, and corresponding updates to the MSA's allocation and the Terms of Use.

## **10. Multi-Jurisdictional Positioning and Framework Alignment**

This Policy is designed principally for SB 26-189 and FERPA and to support Customer compliance with analogous frameworks. AI-specific, privacy, and consumer-protection laws are evolving across jurisdictions; this Policy is read and applied consistent with the principles below.



## **10.1 United States State AI Laws**

Customers deploying the AI Agent for End Users outside Colorado are responsible for identifying and complying with applicable state AI laws, including the Texas Responsible Artificial Intelligence Governance Act (intent-based prohibited practices, with a NIST AI RMF affirmative defense), the California automated-decisionmaking-technology regulations (which turn on technology that substantially replaces human decisionmaking for significant decisions, including school admission), the Utah Artificial Intelligence Policy Act, California SB 942, California AB 2013, the family of state student-data-privacy laws modeled on California's Student Online Personal Information Protection Act (SOPIPA) (which restrict using student data for non-educational purposes, profiling, targeted advertising, and sale), and analogous legislation in other states. The prohibited uses in Section 3 are drafted to be broadly consistent with these regimes but do not enumerate every requirement. The human-in-the-loop design is the most valuable single fact across these frameworks.

## **10.2 Federal Law**

There is no comprehensive federal AI law as of the date of this Policy. Federal sector-specific requirements continue to apply, including as administered by the Federal Trade Commission, the Consumer Financial Protection Bureau, the Equal Employment Opportunity Commission, the Department of Health and Human Services, and the Department of Education. the Children's Online Privacy Protection Act (COPPA) (15 U.S.C. 6501 et seq.; 16 C.F.R. Part 312) applies to services directed to children under 13 or where the operator has actual knowledge of under-13 users; a higher education CXP not directed to children is better positioned, with residual risk in actual knowledge of under-13 users. This Policy is designed to be consistent with the NIST AI RMF as a voluntary federal framework.

## **10.3 European Union (Treated as Imminent)**

With European Union and United Kingdom deployment now near-term, the European analysis is treated as operative. The applicable timing is staged: the Article 50 transparency obligations apply from August 2, 2026, while under the Digital Omnibus political agreement of May 7, 2026 the Annex III standalone high-risk regime (which includes education and vocational training) is deferred to December 2, 2027. That agreement is provisional and takes legal effect on formal adoption, expected before August 2, 2026; until then the original calendar applies. Customers deploying for End Users in the European Union are responsible for compliance with Regulation (EU) 2024/1689 (the EU AI Act), including the Article 50 transparency obligations and, where applicable, the high-risk obligations in Chapter III as they apply to systems within Annex III, Point 3 (education and vocational training). The AI disclosure provisions in Section 5 are drafted to be consistent with Article 50. The Annex III classification analysis is conducted separately for any European Union deployment and does not rely on the ADMT Information-Tool Exclusion, which is a creature of Colorado law. Under Article 5 of the EU AI Act, inferring emotions in education settings is restricted, tied to biometric data, with a narrow exception for medical or safety reasons. Pathify's position is that the Agent's text-based detection is likely outside that restriction because it analyzes text rather than a biometric signal, and that voice, being transcription-only (Section 4.6), sits in the same position as typed text; this holds only while voice remains transcription-only. The planned live-voice capability described in Section 4.6,

which would process the voice signal itself, is treated as a distinct capability that would process biometric data and is subject to separate analysis under Article 5; it may be excluded from European Union and United Kingdom deployments. The General Data Protection Regulation applies in parallel, including Article 22 on automated decisions and Article 27 on European Union representatives, for which Pathify is preparing the corresponding measures in step with its near-term European Union and United Kingdom deployment.

#### **10.4 Other International**

Customers deploying outside the United States and the European Union are responsible for applicable AI, privacy, and consumer-protection laws in those jurisdictions, including the United Kingdom Data (Use and Access) Act 2025 and Information Commissioner's Office guidance, the Australian Privacy Act 1988 (Cth), the New Zealand Privacy Act 2020, Quebec Law 25 automated-decisionmaking provisions, and Canadian federal AI legislation as enacted. Jurisdiction-specific drafting provided by local counsel is integrated as it is finalized.

#### **10.5 No Warranty of Sufficiency**

Pathify provides reasonable assistance with Customer compliance but does not warrant that this Policy alone is sufficient for all laws applicable to a Customer's deployment. Customers retain primary responsibility for identifying applicable laws and implementing deployer-side obligations.

#### **10.6 Standards and Framework Alignment**

Pathify maintains substantive alignment of the AI Agent and this Policy with recognized frameworks as they apply in higher education, used as a structured way to explain governance over time rather than a representation that any framework has been fully adopted in every respect: the NIST AI RMF (the risk-management provisions of Sections 7, 8, and 9 and the role allocation in Section 8 support the Govern, Map, Measure, and Manage functions); ISO/IEC 42001 (Section 9.2 identifies it as a reasonableness benchmark); the FERPA school official framework (Section 4.3); the UNESCO Guidance for Generative AI in Education and Research (2023); the OECD AI Principles; and EDUCAUSE guidance. This Section does not warrant compliance with any framework or commit Pathify to automatic adoption of revisions; it reflects current alignment and an ongoing diligence commitment. Supporting governance artifacts, including Pathify's AI governance statement and the AI Agent risk register, are available to Customers on request.

#### **10.7 NIST AI RMF Crosswalk**

Govern: Sections 1, 8, 9, and 12 establish governance through role allocation, function-based activation conditions, accountability, review cycles, and documented approval paths for Covered ADMT functions. Map: Sections 1, 2, 4, 5, and 6 map intended uses, prohibited uses, system boundaries, model-provider involvement, data categories and sources, human-review points, and function classification. Measure: Sections 6, 7, and 9 require impact documentation, discrimination-risk assessment, feasible pre-activation testing, monitoring of override patterns and known limitations, and documentation of foreseeable misuse and mitigation. Manage: Sections 3, 8, 9, 11, and 12 manage risk through prohibited-use rules, activation gating,

Independent Determination, developer-and-deployer allocation, incident and escalation procedures, material-change management, and corrective action.

## 10.8 Public Higher-Education AI Framework Crosswalk

In addition to the NIST AI RMF crosswalk in Section 10.7 above, which maps this Policy to a recognized AI risk-management framework, this crosswalk maps the AI-framework requirements that United States public higher-education systems now place on their campuses to the provisions of this Policy that address them. The two are complementary: Section 10.7 addresses alignment with a recognized risk-management framework, while this Section addresses the institutional AI-policy requirements those systems impose. The reference set is drawn from the systemwide AI policies and responsible-AI principles that leading public university systems have adopted, principally the State University of New York (SUNY) Systemwide Artificial Intelligence Policy (adopted April 30, 2026), the University of California Responsible AI Principles, and the University System of Georgia (USG) Artificial Intelligence Guidelines (a companion guide to the USG Information Technology Handbook), together with the City University of New York (CUNY) and California Community Colleges responsible-AI principles. These requirements are representative of what public institutions across multiple states now require, and this crosswalk is not limited to any one system. Requirements imposed by other states where a Customer operates are addressed in Section 10.1, and Sections 10 and 12 are structured so that this Policy can extend to additional state and federal AI requirements as they take effect. This crosswalk is a navigational aid and does not warrant compliance; each Customer retains primary responsibility for meeting its own obligations, as provided in Section 10.5 (No Warranty of Sufficiency). Security and procurement reviewers will also require artifacts that sit outside this Policy; Pathify provides those through the Data Processing Addendum and its security and trust documentation, which may include a System and Organization Controls (SOC) 2 report, a subprocessor list, incident-response and breach-notification terms, data-residency and deletion information, and an accessibility conformance report.

| Public higher-education AI-framework requirement                 | Where this Policy addresses it                                                                                                                                                                                   |
|------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Protect institutional and student data                           | Section 4 (Data Handling), including the FERPA school official architecture (Section 4.3), institution-scoped isolation and no cross-institutional data flows (Section 4.5), and the Data Processing Addendum.   |
| Evaluate AI tools for bias; prevent discriminatory or biased use | Section 3.2 (prohibited discriminatory content) and Section 7 (standing discrimination-risk safeguards and monitoring for discriminatory output; application-level bias testing of the Agent is in development). |
| Preserve human decision-making authority                         | Sections 1.1 and 3.1 (the Agent does not make or materially influence a Consequential Decision), the Independent Determination requirement (Section 1.2), and the staff-decide disclosure (Section 5).           |
| Clarify roles and responsibilities                               | Section 8 (developer and deployer role allocation), the Institution Designee role (Section 1.2), and institutional control (Section 12.1).                                                                       |

|                                                                                                             |                                                                                                                                                                                                                                                      |
|-------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Add procurement safeguards for AI tools that take actions inside institutional systems                      | Section 3.5 (Integrated Action, also called write-back): user-initiated, user-confirmed through a distinct confirmation step, deterministic, and Customer-configured; and the product-surface and activation boundaries in Section 2.7.              |
| Provide transparency and AI disclosure to users                                                             | Sections 5 and 5.1 (AI disclosure and activation acknowledgment) and the model-provider disclosure (Section 2.6).                                                                                                                                    |
| Ensure accessibility of the AI tool                                                                         | Section 4.9 (Web Content Accessibility Guidelines conformance target and accessibility conformance report on request).                                                                                                                               |
| Support security review, incident handling, and subprocessor transparency                                   | Sections 4.8 and 12.3 (operational records and customer security input and escalation), Section 11 (policy enforcement and legal enforcement), Section 12.4 (subprocessor list and procurement support materials), and the Data Processing Addendum. |
| Maintain change management and periodic review                                                              | Section 12 (review and updates) and Section 12.2 (material AI change categories).                                                                                                                                                                    |
| Support staff and user training and documentation                                                           | Section 12.4 (procurement and governance support materials) and the data-flow documentation in Section 4.7.                                                                                                                                          |
| Align with recognized frameworks and remain ready for other state and future state and federal requirements | Sections 10.6 and 10.7 (NIST AI RMF, ISO/IEC 42001, FERPA, and other framework alignment), Section 10.1 (other United States state AI laws), and Sections 10 and 12 (extension as new requirements take effect).                                     |

## 11. Policy Enforcement and Legal Enforcement

A violation of this Policy by a Customer constitutes a material breach of the MSA. Pathify reserves the right to suspend or terminate AI Features access for any Customer or End User that violates this Policy.

Under C.R.S. 6-1-1706, violations of Part 17 are enforced by the Colorado Attorney General through the Colorado Consumer Protection Act and constitute deceptive trade practices, enforceable exclusively by the Attorney General. Before an enforcement action, the Attorney General must provide a 60-day notice and opportunity to cure where the Attorney General deems a cure possible; the cure opportunity does not apply to knowing or repeated violations, and the right-to-cure provision is repealed effective January 1, 2030. Under C.R.S. 6-1-1709, Part 17 creates no new private right of action. Separately, C.R.S. 6-1-1707 governs how fault for unlawful discrimination is allocated between developers and deployers in civil actions under existing law and provides, at C.R.S. 6-1-1707(7), that a contract provision purporting to indemnify a developer or deployer against liability for its own acts or omissions related to the use of automated decision-making technology in making a Consequential Decision in violation of Colorado anti-discrimination law is contrary to public policy and void. SB 26-189 does not displace liability that may arise under other applicable anti-discrimination law.

## 12. Review and Updates

This Policy shall be reviewed at least annually and updated as necessary to reflect changes in the AI Agent's functionality, data flows, or model provider; changes in Applicable Laws, including

rules the Colorado Attorney General is directed to adopt on or before January 1, 2027 under C.R.S. 6-1-1704(4) and 6-1-1705(3); findings from ongoing monitoring; and changes in Customer use cases or configurations that affect a function's regulatory classification. Pathify shall make updated versions available at the URL specified in the MSA definition of "Acceptable AI Policy" and shall provide at least 30 days' prior written notice of material changes that would materially affect a Customer's obligations or a function's classification. Non-material updates may be made without prior notice. Legal characterizations in this Policy are based on the laws and guidance available as of the Policy date and may be updated as rulemaking, codification, and implementation guidance evolve.

### **12.1 Institutional Tool and Institutional Control**

This Policy reflects a single organizing principle: the AI Agent is an institutional tool that operates under the Customer's direction and control. It informs, navigates, refers, answers questions, classifies for human attention, and offers informational suggestions; it does not decide, and it does not provide student-specific recommendations that function as a factor in a Consequential Decision through its Core Question-and-Answer Function. The institution, acting through its qualified staff, remains the decision-maker for every Consequential Decision. The function classifications, the Independent Determination requirement, the prohibitions in Section 3, the disclosures in Section 5, and the data-handling provisions in Section 4 are the mechanisms that keep that principle true in operation. Where future capabilities would change it, this Policy requires the corresponding notice, approval, and safeguards before they take effect.

### **12.2 Material AI Change Categories**

A material AI change includes, at a minimum: (i) a change in the underlying model provider or a materially different model surface; (ii) a change in whether institution-sourced FERPA Records or other personally identifiable information are transmitted to the model provider; (iii) a change in the classification of an AI Feature; (iv) a change that materially alters customer-facing AI functionality or the role of human review; (v) a change in training, retention, regional-processing, or materially relevant upstream-provider terms that Pathify reasonably determines affects the representations in this Policy; or (vi) activation of a new safety, monitoring, evaluation, or classification function not previously in scope for the Customer. Pathify may satisfy notice obligations through this Policy, the DPA, release notes, an AI Feature Authorization Schedule, or other written notice.

### **12.3 Customer Security Input and Escalation**

Pathify maintains a documented process for receiving and triaging good-faith customer-reported security concerns relating to AI Features, evaluates them based on severity, exploitability, and operational impact, and provides reasonable status updates for validated material issues, retaining discretion over remediation method, prioritization, and timing subject to the governing agreement and applicable law.

### **12.4 Procurement Support Materials**

Pathify may support Customer procurement, security, privacy, audit, and AI-governance review through materials that may include a Product Surface and Data-Flow Schedule, an AI Feature

Authorization Schedule, a Subprocessor List, a security input and escalation description, testing or monitoring summaries subject to confidentiality, and other AI governance documentation. Such materials are descriptive unless expressly incorporated into the governing agreement. To support Customers' training and tool-inventory obligations, Pathify will also provide administrator-facing documentation that a Customer may use in its own AI-literacy and staff-training materials, and the AI Feature Authorization Schedule identifies the AI Features a Customer has enabled so the Customer can maintain an inventory of approved AI tools.

### **12.5 Continuity and Institutional Planning**

Pathify designs AI Features to support institutional operations, but the Customer remains responsible for its own continuity planning. Nothing in this Policy guarantees uninterrupted availability of AI Features or continuity of Customer operations during a broader platform, integration, or upstream-provider disruption.

## **13. Contact**

Questions about this Policy should be directed to:

Path Education Inc.  
5910 S University Blvd, C-18 #224  
Greenwood Village, CO 80121  
Tel: +1 (720) 410-2177

---