

GOVERNANCE FRAMEWORK

Template for Texas School Districts

Field	Entry
Aligned With	Texas SB 1964 (eff. Sep 1, 2025) TAC 219 HB 3512 (eff. Sep 1, 2025) HB 149 / TRAIGA SB 12
District	Tornillo Independent School District
Superintendent	Dr. Rosy Vega-Barrio
Technology Director / AI Risk Officer	Carlos Garcia & Christopher Escarsega
Adoption Date	July 20 2026
Next Review	Annually; Tier 3 (HSAI) systems reviewed semi-annually.

SECTION 1: PURPOSE & LEGAL BASIS

1.1 Purpose of This Framework

This Artificial Intelligence Governance Framework establishes **Tornillo Independent School District** policies and procedures for the responsible procurement, deployment, oversight, and evaluation of artificial intelligence systems. It is designed to protect students, staff, and community members while enabling the district to leverage AI tools that enhance educational outcomes.

1.2 Legal Authority

This framework is adopted in compliance with the following:

Governing Law	Reference / Summary
Texas Senate Bill 1964 (89th Leg., eff. Sep 1, 2025)	AI Code of Ethics and minimum governance standards for local governments, including school districts
Texas Government Code §2054	DIR oversight authority
NIST AI Risk Management Framework (AI RMF 1.0)	Federal alignment standard referenced in SB 1964
FERPA (20 U.S.C. §1232g)	Student education records privacy
COPPA	Children's Online Privacy Protection Act
HB 149 / TRAIGA (eff. Jan 1, 2026)	Texas Responsible AI Governance Act; public-facing AI disclosure and biometric identifier requirements
HB 3512 (eff. Sep 1, 2025)	Annual cybersecurity and AI training required for cybersecurity coordinator and any elected/appointed officials who have access to a

Governing Law	Reference / Summary
	local government computer system and utilize it to perform at least 25% of their duties
SB 12	Parental consent required for biometric data storage and health/wellness AI capturing psychological student data
DIR Prohibited Technologies List	All AI tools must be screened against DIR's prohibited software/application/developer list before procurement
TAC 219: DIR AI Code of Ethics	Minimum AI risk management standards, AI Code of Ethics, AI Risk Officer , and heightened scrutiny AI system (HSAI) requirements

1.3 What [TAC 219](#) Requires of School Districts

Under TAC 219 (adopted pursuant to SB 1964), school districts, as local governmental entities, are required to:

- Adopt the AI Code of Ethics established by DIR under TAC §219.11(a)
- Designate an employee as the [AI Risk Officer](#) (TAC §219.21(a)) (this may be an existing employee with added duties)
- Establish a process to identify and inventory all heightened scrutiny AI (HSAI) systems (TAC §219.21(b))
- Conduct a written AI risk assessment before deploying any HSAI system (TAC §219.22(a)-(b))
- Have the designated AI Risk Officer review the completed risk assessment and approve or deny HSAI deployment, then notify the executive head (TAC §219.22(c))
- Retain written risk assessment records per the state records retention schedule (TAC §219.22(d))
- Adopt an acceptable use policy (AUP) for each HSAI system and train all employees on it (TAC §219.24(b))
- Provide risk-specific training to employees who access, use, or manage an HSAI system (TAC §219.24(c))
- Contractually require vendors deploying HSAI systems to implement the [NIST AI Risk Management Framework](#) (TAC §219.24(d))
- Consider conducting a full AI impact assessment when deploying an HSAI system, in alignment with state agency requirements (TAC §219.23(e) — encouraged, not mandatory for school districts)
- Comply with all DIR requirements regarding procurement of information resources that include an HSAI system (TAC §219.23(e)(1))

1.4 Scope

This framework applies to all AI systems procured, developed, or deployed by Tornillo Independent School District, including:

- Instructional AI tools and adaptive learning platforms
- Administrative AI systems (student information, HR, finance)
- Security and surveillance AI (facial recognition, behavior monitoring)
- Communication AI tools (chatbots, automated messaging)
- Vendor-supplied AI embedded in purchased software

SECTION 2: KEY DEFINITIONS

The following definitions are aligned with Texas Government Code §2054 as amended by SB 1964 and TAC 219:

Term	SB 1964 / TAC 219 / District Definition
Artificial Intelligence (AI) System	A machine-based system that, for a given set of objectives, makes predictions, recommendations, or decisions influencing real or virtual environments. (Tex. Gov. Code §2054.003(1-a))
Heightened Scrutiny AI System (HSAI)	An AI system that serves as a controlling factor or principal basis in decisions affecting an individual's rights, benefits, privileges, or obligations — requiring additional oversight, documentation, and bias testing.
Controlling Factor	An AI system's output that singularly determines or substantially drives a consequential outcome without meaningful human review. (Tex. Gov. Code §2054.003(2-c))
Principal Basis	An AI system's output that serves as the primary rationale or driver behind a decision, even if human review technically occurs. (Tex. Gov. Code §2054.003(11))
Consequential Decision	A decision that materially affects an individual's rights, benefits, privileges, or obligations. (Tex. Gov. Code §2054.003(2-b))
AI Code of Ethics	The statewide code established by DIR under TAC §219.11(a) setting minimum standards for AI fairness, transparency, accountability, and human oversight. School districts are required to adopt this code.
AI Risk Officer	A designated district employee responsible for promoting ethical AI procurement, ensuring risk assessments are completed, reviewing HSAI deployments, and notifying the executive head of deployment decisions. (TAC §219.21) — May be an existing employee with added duties.
Deployer	The district or department that uses or operates an AI system — distinct from the developer who built it.
AI System Inventory	A documented registry of all HSAI systems in active use by the district, maintained and provided to DIR upon request. (TAC §219.21(b))
Biometric Data	Data derived from unique biological characteristics such as faceprints, voiceprints, or fingerprints used to identify an individual. Requires prior written parental consent before collection or storage for students. (TEC 26.009; SB 12)
Disablement Protocol	A documented technical or operational procedure to immediately disable an AI system if harmful, inaccurate, or unauthorized outputs are detected. (TAC §219.11(c)(2)(C))
DIR-Certified AI Training	Annual AI literacy training programs certified by DIR as required by HB 3512, covering AI fundamentals, ethics, bias awareness, and responsible use in a government/education context. (Applies to cybersecurity coordinator AND certain elected/appointed officials with ≥25% computer)
Personal Identifying Information (PII)	As defined by Texas Business & Commerce Code §521.002(a)(1).
Unlawful Harm	As defined by Texas Government Code §2054.701.

SECTION 3: AI SYSTEM RISK CLASSIFICATION

SB 1964 and TAC 219 establish a tiered approach to AI oversight. Districts must classify AI systems before deployment and apply the corresponding level of scrutiny and controls.

3.1 Risk Tiers

Tier	Description	District Examples
TIER 1 Low Risk	AI that assists humans with no direct effect on individual rights, benefits, or privileges.	Scheduling assistants, curriculum recommendation tools, grammar checkers, AI chatbots for general information.
TIER 2 Moderate Risk	AI that informs decisions affecting individuals but is not the controlling factor. Human review is meaningful and documented.	Student early-warning/intervention systems, attendance pattern analysis, adaptive assessment platforms, AI tutoring tools with intervention data.
TIER 3 High Risk (HSAI)	AI serving as a controlling factor or principal basis in decisions affecting rights, benefits, or privileges. Requires full heightened scrutiny protocol under TAC §219.22 and §219.24.	Automated disciplinary scoring, special education eligibility determination, employee performance AI, facial recognition for access control, health/wellness AI capturing psychological data for evaluation.

3.2 Classification Decision Process

Before deploying any AI system, the AI Risk Officer will complete the AI System Classification Worksheet (Appendix A) and obtain approval from the AI Team. The following questions guide classification:

- Does the AI system's output directly determine an outcome affecting an individual's rights, benefits, or enrollment without substantive human review? → If YES: Tier 3 (HSAI)
- Does the AI system's output serve as the "principal basis" for a human decision about an individual, even if a human technically makes the final call? → If YES: Tier 3 (HSAI)
- Does the AI system's output inform a human decision-maker who has genuine authority to override the system? → If YES: Tier 2
- Is the AI system purely assistive with no individual-level decision output? → If YES: Tier 1
- Does the AI system collect biometric data (faceprints, voiceprints)? → If YES: Biometric Consent Protocol required (Section 12) regardless of tier
- Is this a health/wellness AI capturing psychological student data? → If YES: Tier 3 + Parental Consent required (Section 12.2)

SECTION 4: AI SYSTEM INVENTORY

TAC §219.21(b) requires school districts, as local governmental entities, to establish a process to identify and inventory all HSAI systems and provide this information to DIR upon request.

4.1 District AI System Inventory Register

System Name / Vendor	Department / Campus	Risk Tier	Date Deployed	Date of Last Review	Brief Description of AI Function	Responsible Party
Google Gemini for Education	District-Wide	Tier 1 (Low Risk)	July 22, 2025	June 22, 2026	Generative AI platform used for instructional support, lesson planning, research assistance, content creation, and	Technology Director / AI Risk Officer (Christopher Escarsega)

System Name / Vendor	Department / Campus	Risk Tier	Date Deployed	Date of Last Review	Brief Description of AI Function	Responsible Party
					administrative productivity.	
Microsoft 365 Copilot	District-Wide	Tier 1 (Low Risk)	July 22, 2025	June 22, 2026	AI-powered productivity assistant integrated with Microsoft 365 for document creation, email drafting, meeting summaries, data analysis, and workflow automation.	Technology Director / AI Risk Officer (Christopher Escarsega)
ChatGPT / OpenAI	District-Wide	Tier 1 (Low Risk)	July 22, 2025	June 22, 2026	Generative AI platform used by staff and teachers for lesson planning, professional communications, instructional content development, research assistance, and administrative support.	Technology Director / AI Risk Officer (Christopher Escarsega)
Canva / Canva Pty Ltd	District-Wide	Tier 1 (Low Risk)	July 22, 2025	June 22, 2026	AI-powered graphic design and content creation suite (Canva Magic Studio) used by staff and students for generating presentations, lesson materials, image/video editing, text-to-image generation, and collaborative visual projects.	Technology Director / AI Risk Officer (Christopher Escarsega)

Note: The district has approved Google Gemini for Education, Microsoft 365 Copilot, and ChatGPT for authorized educational and administrative use. Staff must comply with district data privacy requirements and shall not enter confidential student information, personally identifiable information (PII), special education records, discipline records, medical information, or other protected district data into AI systems unless specifically authorized and protected by an approved data privacy agreement.

Staff shall only use district-approved enterprise versions of AI tools that provide contractual privacy protections. Public consumer AI tools may not be used for student records, employee records, confidential district information, or protected educational data unless specifically approved by the Technology Department.

Students may use district-approved AI tools only when permitted by the classroom teacher and in accordance with academic integrity expectations.

Responsible Party: Technology Director / AI Risk Officer (Christopher Escarsega)

SECTION 5: AI CODE OF ETHICS

TAC §219.11(a) requires every school district to adopt the AI Code of Ethics established by DIR and follow the ethical principles in the procurement, development, deployment, or use of AI systems. This section establishes Tornillo Independent School District ethical commitments, aligned with SB 1964, TAC 219, and the [NIST AI Risk Management Framework](#).

Principle (TAC §219.11)	District Commitment
Human Oversight & Control(§219.11(c))	All AI systems used by the district shall be subject to meaningful human oversight. AI systems must be deployed in ways that enable staff to review and analyze inputs and outputs throughout the AI lifecycle. HSAI systems require heightened human oversight. No AI system shall make a final, unreviewed decision affecting a student's or employee's rights, benefits, or enrollment status. AI systems must be capable of being disabled until harmful or inaccurate outputs can be remedied.
Fairness(§219.11(d))	The district shall ensure its use of AI systems does not infringe upon the legally protected rights and liberties of students and staff or result in unlawful harm. Data governance practices for AI systems shall be implemented throughout the AI system's lifecycle to ensure fairness. Disparate impact analysis shall be conducted for any Tier 3 system.
Accuracy(§219.11(e))	Staff shall be trained to understand the importance of verifying AI outcomes for accuracy. Processes shall be formalized for monitoring system accuracy before deployment and throughout the AI lifecycle, as accuracy may change over time.
Redress(§219.11(f))	A mechanism shall be provided to seek redress when an AI system makes a consequential decision that unfairly impacts a student, employee, or other individual in a material way. A designated point of contact shall handle such inquiries. Internal procedures shall allow staff to identify and remedy negative impacts caused by AI systems.
Transparency(§219.11(g))	The district shall disclose when individuals interact with an AI system and when an AI system is used to make material decisions about their rights or access to district services. AI systems shall never be represented as human when interacting with students, families, or the public. All public-facing AI must identify itself as AI before interacting with users (HB 149 / BCC §552.051).
Data Privacy & Data Minimization(§219.11(h))	The district shall collect and maintain only the PII needed for operations and establish a process to delete PII consistent with records retention schedules. Employees shall be trained about the risk of inputting sensitive data or PII into publicly available AI systems. AI systems shall comply with FERPA, COPPA, and Texas data privacy law. Student data shall not be used to train external AI models without explicit board approval and parental consent.
Security(§219.11(i))	AI systems shall be monitored, secured, and tested to prevent or limit security attacks. AI system providers must disclose known vulnerabilities and resolutions in a timely manner. AI systems must meet NIST SP 800-53 security controls.
Accountability & Liability(§219.11(j))	Training shall be provided to employees on how to use AI systems effectively, safely, and ethically. Vendors shall be contractually bound to these AI ethical principles and applicable laws. All AI systems deployed must comply with state and federal law. Appropriate retention schedules shall be established for AI system records, including Public Information Act implications.
Evaluation(§219.11(k))	Methods shall be established for regular evaluation of AI systems to ensure ongoing educational benefit. All AI systems shall be evaluated annually; HSAI systems shall be reviewed semi-annually. Evaluations shall be documented.
Documentation(§219.11(l))	Records shall be maintained of the data sources used in AI systems and how AI systems are modified throughout their lifecycle. Documentation of preliminary assessments, ongoing monitoring, and complaints shall be maintained throughout each AI system's lifecycle.
Prohibited Technologies	The district shall screen all AI tools against the DIR Prohibited Technologies List and the Texas Governor's Executive Order prohibiting social media applications on government devices/networks before procurement or deployment.

Principle (TAC §219.11)	District Commitment
Biometric Protection(SB 12 / TEC 26.009)	The district shall not deploy AI systems that uniquely identify students using biometric data without prior written parental consent. No AI system shall collect, store, or use student biometric data, including faceprints, voiceprints, or fingerprints, without documented prior written consent in the student's file.

SECTION 6: PROCUREMENT & APPROVAL PROCESS

6.1 Approval Workflow by Risk Tier

Step	Tier 1 (Low)	Tier 2 (Moderate)	Tier 3 (High / HSAI)
AI System Classification Worksheet	Required	Required	Required
Technology Director Review	Required	Required	Required
DIR Prohibited Technologies Screening	Required	Required	Required
DIR Prohibited Technologies List	Required	Required	Required
Data Privacy Impact Assessment	Optional	Required	Required
Bias & Equity Impact Assessment	Not Required	Required	Required (with demographic audit)
Cybersecurity Review (NIST controls)	Basic	Standard	Full NIST SP 800-53 review
Vendor AI Compliance Attestation	Recommended	Required	Required
Vendor NIST AI RMF Contractual Requirement(TAC §219.24(d))	Not Required	Encouraged	Required — vendor must contractually implement NIST AI RMF
Written AI Risk Assessment (TAC §219.22)	Not Required	Not Required	Required before deployment or material change
HSAI-Specific Acceptable Use Policy (AUP)(TAC §219.24(b))	Not Required	Not Required	Required — all employees using the system must be trained on it
Risk-Specific Employee Training(TAC §219.24(c))	Not Required	Not Required	Required for all employees who access, use, or manage the HSAI
Superintendent Approval	Not Required	Suggested	Suggested
Board Notification	Not Required	Informational	Formal Board Approval Suggested
Biometric Consent Protocol (if applicable)	As needed	As needed	Required if biometric data involved
Public Website Disclosure	Optional	Suggested	Required
HSAI Risk Assessment Record Retention(TAC §219.22(d))	Not Required	Not Required	Required (retain written risk assessment per state records schedule)

Step	Tier 1 (Low)	Tier 2 (Moderate)	Tier 3 (High / HSAI)
AI Impact Assessment (TAC §219.23(e))	Not Required	Not Required	Strongly encouraged (not mandatory for districts, but considered best practice)

6.2 Vendor Contract Requirements

All contracts with AI system vendors must include provisions requiring the vendor to:

- Disclose whether and how AI is used in the product, including embedded AI features
- Disclose which AI models underlie the product (including any third-party models) to support [DIR Prohibited Technologies](#) and Governor's Executive Order compliance
- Provide documentation of bias testing, training data sources, and model performance metrics
- Comply with FERPA, COPPA, and Texas student data privacy law
- Agree not to use district data to train AI models without written consent
- Allow third-party security audits upon district request
- Notify the district within 48 hours of any data breach involving district data
- For HSAI (Tier 3): contractually implement the [NIST AI Risk Management Framework](#) (TAC §219.24(d))
- Be bound to the district's AI ethical principles and applicable laws (TAC §219.11(j)(2)(B))
- Allow student-level data export to fulfill FERPA educational record requests, including AI-generated intervention or tutoring records (SB 12; TEC §26.004(b))

SECTION 7: AI RISK ASSESSMENT FORM (TAC §219.22)

[TAC §219.22](#) requires school districts to conduct a written AI risk assessment before deploying any HSAI (Tier 3) system and at the time any material change is made to the system. Complete this form for all Tier 3 systems before deployment. **Districts are also encouraged to complete this form for Tier 2 systems as a best practice.**

 **Note:** [TAC §219.23](#) (Full AI Impact Assessment) is expressly required for state agencies. Under TAC §219.23(e), school districts are required to comply with DIR procurement requirements for HSAI systems and are encouraged, but not mandated, to conduct a full impact assessment. Districts that choose to complete the full impact assessment should use Parts A–E of this form.

Part A: System Identification

Field	Response
System Name:	Gemini
Vendor / Developer:	Google
Department / Campus:	District-Wide
Primary Contact:	
Risk Tier Classification:	Tier 1 (Low Risk)
Assessment Date:	July 22, 2025
Description of AI System Purpose and Function:	Core Purpose: A native AI collaborator embedded directly across the Google Workspace for Education suite (including

Field	Response
	Docs, Sheets, Slides, Drive, and Gmail) designed to optimize teacher workflows and student productivity. Data Privacy & Security: Enforces enterprise-grade data protections mapped to the school's existing Google Admin security settings. The environment is entirely ad-free. All user interactions stay strictly within the institution's domain and are never human-reviewed or used to train public Google AI models.
Describe the decisions or outputs the AI system makes or influences:	Key Outputs: Provides inline writing and editing help in Docs, generates data-driven formulas and categorizations in Sheets, creates slide outlines and visual graphics, and summarizes long text strings from Drive or Gmail. Human Control: Gemini acts solely as a supportive drafting partner. It holds zero autonomous authority over administrative operations or student evaluations. All lesson structures, data models, and campus communications must be reviewed and authorized by a human operator.
Does this system collect or process biometric data (faceprints, voiceprints)? If YES, is prior written parental consent in place? (TEC 26.009; SB 12)	☐ Yes — consent mechanism: _____ ☒ No biometric data
Is this a health/wellness tool that captures psychological student data for evaluation or treatment? (TEC 26.009; SB 12)	☐ Yes — parental consent required ☒ No
Has the vendor been screened against the DIR Prohibited Technologies List ?	☒ Yes — cleared ☐ No — action required ☐ Pending

Field	Response
System Name:	Copilot
Vendor / Developer:	Microsoft 365
Department / Campus:	District-Wide
Primary Contact:	
Risk Tier Classification:	Tier 1 (Low Risk)
Assessment Date:	July 22, 2025
Description of AI System Purpose and Function:	Core Purpose: A secure AI assistant used by faculty, staff, and students to streamline lesson planning, accelerate research, and automate administrative tasks. Data Privacy & Security: Protected by Enterprise Data Protection (EDP) . All data and prompts are strictly encrypted and kept private within the school's network. They are never used to train public AI models.
Describe the decisions or outputs the AI system makes or influences:	Key Outputs: Generates first drafts of educational materials (quizzes, rubrics, lesson plans), summarizes long documents/meetings, and analyzes spreadsheet data.

Field	Response
	Human Control: The AI cannot make autonomous decisions, grade assignments, or enact policy. It acts purely as a drafting assistant, and all outputs require human review and approval.
Does this system collect or process biometric data (faceprints, voiceprints)? If YES, is prior written parental consent in place? (TEC 26.009; SB 12)	☐ Yes — consent mechanism: _____ ☒ No biometric data
Is this a health/wellness tool that captures psychological student data for evaluation or treatment? (TEC 26.009; SB 12)	☐ Yes — parental consent required ☒ No
Has the vendor been screened against the DIR Prohibited Technologies List ?	☒ Yes — cleared ☐ No — action required ☐ Pending

Field	Response
System Name:	ChatGPT
Vendor / Developer:	OpenAI
Department / Campus:	District-Wide
Primary Contact:	
Risk Tier Classification:	Tier 1 (Low Risk)
Assessment Date:	July 22, 2025
Description of AI System Purpose and Function:	Core Purpose: An AI assistant used by faculty, staff, or students to assist with drafting text, brainstorming creative instructional ideas, and formatting open-source information. Data Privacy & Security: Note for Institutional Compliance: Because this is a standard account (not ChatGPT Edu), data is <i>not</i> shielded by default enterprise agreements. User prompts and uploaded files may be used by OpenAI to train public models. To maintain institutional privacy, users must manually go to Settings > Data Controls and turn off "Improve the model for everyone". No confidential institutional data or personally identifiable student data should be entered.
Describe the decisions or outputs the AI system makes or influences:	Key Outputs: Generates text drafts, structures basic lesson frameworks, provides coding/scripting assistance, and formats public datasets. Human Control: Acts strictly as an external, consumer-grade drafting aid. It has zero integration with school systems and holds no autonomous authority over grading, evaluations, or policy. All outputs require mandatory human verification.
Does this system collect or process biometric data (faceprints, voiceprints)? If YES, is prior written parental consent in place? (TEC 26.009; SB 12)	☐ Yes — consent mechanism: _____ ☒ No biometric data
Is this a health/wellness tool that captures psychological student data for	☐ Yes — parental consent required ☒ No

Field	Response
evaluation or treatment? (TEC 26.009; SB 12)	
Has the vendor been screened against the DIR Prohibited Technologies List ?	☒ Yes — cleared ☐ No — action required ☐ Pending

Field	Response
System Name:	Canva / Canva Pty Ltd
Vendor / Developer:	OpenAI
Department / Campus:	District-Wide
Primary Contact:	
Risk Tier Classification:	Tier 1 (Low Risk)
Assessment Date:	July 22, 2025
Description of AI System Purpose and Function:	Core Purpose An AI-powered graphic design and multimedia creation suite (via Canva Magic Studio) used by faculty, staff, and students to assist with drafting text (Magic Write), generating visual elements, building interactive presentations, organizing digital whiteboards, and formatting instructional templates. Note for Institutional Compliance: Under our Canva for Education institutional framework, student and staff data is rigorously protected. Canva is a signatory of the National Data Privacy Agreement (NDPA) and is fully FERPA and COPPA compliant. Through the <i>Canva Shield</i> security umbrella, Canva explicitly states that user content, prompts, and uploaded assets from Education accounts are not used to train public or proprietary AI models. Furthermore, the Technology Director retains absolute administrative control to toggle or restrict specific Magic Studio AI tools by user role or campus.
Describe the decisions or outputs the AI system makes or influences:	Key Outputs Generates custom graphic layouts, presentation decks, text-to-image/video media assets, text modifications (summarizing, rewriting, tone adjustment), and automated classroom activities or quiz frameworks. Human Control Acts strictly as a collaborative visual and textual drafting aid. It operates completely independently of core student information systems (SIS) and holds no autonomous authority over grading, academic evaluations, or district policy. All AI-generated designs and textual outputs require mandatory human review and verification before classroom deployment or grading.
Does this system collect or process biometric data (faceprints, voiceprints)? If YES, is prior written parental consent in place? (TEC 26.009; SB 12)	☐ Yes — consent mechanism: _____ ☒ No biometric data

Field	Response
Is this a health/wellness tool that captures psychological student data for evaluation or treatment? (TEC 26.009; SB 12)	☐ Yes — parental consent required ☒ No
Has the vendor been screened against the DIR Prohibited Technologies List ?	☒ Yes — cleared ☐ No — action required ☐ Pending

Part B: Security & Performance Risk Assessment (TAC §219.22(b))

This section directly addresses the written risk assessment requirements of TAC §219.22(b) for HSAI systems.

TAC §219.22(b) Requirement	Google Gemini: Response / Evidence
Known security risks of this HSAI system (§219.22(b)(1))	Malicious Document Injection: Attackers can craft files with hidden instructions that, when analyzed by Gemini via Google Drive integrations, execute malicious commands. Cross-Tenant Logic Failures: Probabilistic confusion if multiple documents with conflicting internal operational logic are analyzed concurrently. Workspace Access Creep: Risks stemming from overly broad file sharing permissions within Google Drive, allowing Gemini to pull data from files that users should not access.
Mitigation steps available to limit identified security risks (§219.22(b)(1))	Admin Console Control: Google Workspace Administrators must centrally manage Gemini settings, restricting or allowing access strictly by Organizational Units (OUs). Drive Permissiveness Audit: Continuous automated auditing of Google Drive file-sharing parameters (internal vs. public links) to limit the context window scope of Gemini. Prompt Sanitization Training: Educating institutional staff on checking and stripping out any hidden metadata or structural strings before analyzing external public documents.
Accuracy of system outputs — measurements of accuracy and relevance (§219.22(b)(2)(A))	Google does not provide quantified accuracy logs directly to the institution. Relevance is structurally enhanced through Google Workspace Extensions and direct Google Search grounding. Measurement Standard: Enforces mandatory Human-in-the-Loop (HITL) evaluation. Users are instructed to verify any critical data points using the "Double-Check Response" tool, which runs secondary Google searches to cross-examine claims.
Operational performance metrics — model latency, uptime, and error rate (§219.22(b)(2)(B))	Latency: Fast processing cycles, averaging 1 to 4 seconds for native text generation across Docs, Slides, and Gmail.

TAC §219.22(b) Requirement	Google Gemini: Response / Evidence
	Uptime: Protected under the core Google Workspace for Education Service Level Agreement (SLA), aiming for a 99.9% operational threshold. Error Rate: Low system-level error rates; exceptions and localized generation failures are tracked through Google Workspace Admin status reporting dashboards.
Description of the system's algorithms and how it makes decisions (§219.22(b)(3)(A))	Powered by Google's native multimodal Gemini deep learning models, built on highly optimized transformer architectures. Decision-Making Mechanism: Operates purely through neural-network probabilistic weights. It processes input instructions across text, code, images, or audio concurrently to output the statistically most fitting data generation matrix, holding no actual administrative logic or self-determined intent.
Data used to train the system's model (§219.22(b)(3)(B))	Pre-trained on Google's extensive, multi-billion-node multimodal web indexes, open source archives, literature, and specifically compiled training sets. Crucial Compliance Note: Under the Google Workspace Education agreement, all customer data, prompts, and created assets are stored privately within the school's domain and are never harvested for public model training or reviewed by human data annotators at Google.
Availability of inputs and outputs for ongoing monitoring of decision-making (§219.22(b)(3)(C))	Available via the Google Admin Console. System administrators can pull comprehensive usage reports, audit logs, and access records to monitor deployment patterns across the institutional ecosystem.

TAC §219.22(b) Requirement	Microsoft Copilot: Response / Evidence
Known security risks of this HSAI system (§219.22(b)(1))	Indirect Prompt Injection: Malicious commands hidden inside external documents or web pages processed by Copilot could hijack the model's behavior. Over-Privileged Access (Internal Leakage): If a school's internal SharePoint or OneDrive permissions are misconfigured by IT, Copilot may accidentally surface confidential files to unauthorized users within the institution. System Prompt Leakage: Attackers using sophisticated prompt-chaining techniques could extract internal organizational guardrails or system rules.

TAC §219.22(b) Requirement	Microsoft Copilot: Response / Evidence
Mitigation steps available to limit identified security risks (§219.22(b)(1))	Data Governance Enforcement: Strictly enforce the principle of least privilege across institutional SharePoint and OneDrive environments before rolling out Copilot. Enterprise Data Protection (EDP): Enforce mandatory login using the institutional tenant, ensuring all prompts are encrypted and completely isolated from public data streams. User Training: Mandatory training sessions instructing staff and students to never treat Copilot as a source of secure authentication or data boundary controls.
Accuracy of system outputs — measurements of accuracy and relevance (§219.22(b)(2)(A))	There is no objective, automated numerical accuracy score provided natively by Microsoft. Relevance is optimized via Retrieval-Augmented Generation (RAG) through the Microsoft Graph. Measurement Standard: Evaluation relies strictly on Human-in-the-Loop (HITL) verification. Users must manually review, fact-check, and cross-reference all outputs against authoritative sources before deployment.
Operational performance metrics — model latency, uptime, and error rate (§219.22(b)(2)(B))	Latency: Average token generation latency spans 2 to 5 seconds depending on query complexity and system load. Uptime: Governed by the broader Microsoft 365 Service Level Agreement (SLA), targeting a 99.9% uptime baseline. Error Rate: Monitored via the Microsoft 365 Service Health Dashboard; standard API timeouts or processing errors are reported through IT tenant alerts.
Description of the system's algorithms and how it makes decisions (§219.22(b)(3)(A))	Copilot runs on advanced generative transformer architectures (primarily OpenAI's GPT models) integrated with Microsoft's proprietary indexing algorithms. Decision-Making Mechanism: The system does not possess deterministic logic or agency. It makes statistical, probabilistic predictions to generate the most relevant sequence of words (tokens) based on the context of the user's prompt combined with the retrieved tenant data.
Data used to train the system's model (§219.22(b)(3)(B))	Pre-trained on immense, publicly available datasets from the internet, digitized books, academic repositories, open-source code, and licensed media. Crucial Compliance Note: No institutional files, prompt histories, or student/faculty interactions from the school's environment are used to train or modify the underlying model.

TAC §219.22(b) Requirement	Microsoft Copilot: Response / Evidence
Availability of inputs and outputs for ongoing monitoring of decision-making (§219.22(b)(3)(C))	Fully available to institutional IT administrators. All user inputs (prompts) and outputs can be centrally tracked, archived, and audited using the Microsoft Purview compliance portal and standard M365 unified audit logs.

TAC §219.22(b) Requirement	ChatGPT Response / Evidence
Known security risks of this HSAI system (§219.22(b)(1))	Public Data Leakage: Because this is a standard consumer-grade account, user prompts and uploaded institutional files are retained by default and may be integrated into public training models. Shadow AI Exposure: Individual usage lacks centralized enterprise IT firewall protections, increasing susceptibility to credential phishing and account-takeover attacks. Model Inversion / Privacy Leakage: Susceptibility to targeted reverse-engineering prompts where malicious actors attempt to extract sensitive data previously input by other users.
Mitigation steps available to limit identified security risks (§219.22(b)(1))	Mandatory Configuration Change: Every user must manually navigate to Settings > Data Controls and turn off "Improve the model for everyone" to stop data training. Strict PII Ban: Complete institutional prohibition of inputting Personally Identifiable Information (PII), FERPA-protected student records, or sensitive financial data. Browser/Network Guardrails: Implementation of endpoint monitoring or web-filtering tools to block unapproved extensions interacting with the ChatGPT web interface.
Accuracy of system outputs — measurements of accuracy and relevance (§219.22(b)(2)(A))	No baseline metrics are provided to institutional users. Accuracy fluctuates dramatically depending on the model version used (free vs. plus tiers) and the specific academic domain being queried. Measurement Standard: Rely entirely on individual, manual user verification. The system exhibits a recognized baseline risk for "hallucinations" (confident fabrication of facts) on ungrounded data.
Operational performance metrics — model latency, uptime, and error rate (§219.22(b)(2)(B))	Latency: Highly variable. Free tiers experience significant degradation during peak traffic hours, while paid tiers average 1 to 3 seconds. Uptime: Subject to public OpenAI status logs. No enterprise-backed availability or uptime guarantees are provided for standard tier accounts.

TAC §219.22(b) Requirement	ChatGPT Response / Evidence
	Error Rate: Subject to standard consumer network drops or API rate-limiting blocks, which are handled individually at the user interface level.
Description of the system's algorithms and how it makes decisions (§219.22(b)(3)(A))	Utilizes OpenAI's proprietary Generative Pre-trained Transformer (GPT) neural network architecture. Decision-Making Mechanism: It relies on deep learning weights to predict text paths autoregressively. It processes information linearly and structurally based on past data correlations, carrying no actual semantic comprehension or authoritative decision-making weight.
Data used to train the system's model (§219.22(b)(3)(B))	Trained on diverse web text corpuses, open-source repositories, articles, books, and user-submitted conversations. Compliance Risk: Unless the "Data Controls" opt-out is manually selected by the end-user, all inputs typed into standard accounts are routinely processed for future model training cycles.
Availability of inputs and outputs for ongoing monitoring of decision-making (§219.22(b)(3)(C))	Extremely Limited: Central IT administration has zero centralized visibility, dashboard tracking, or auditing capabilities for individual standard accounts. Ongoing monitoring is confined exclusively to the local chat history panel visible to the individual user.

TAC §219.22(b) Requirement	Canva Response / Evidence
Known security risks of this HSAI system (§219.22(b)(1))	Centralized Admin Controls (Canva Shield): The Technology Director holds master controls to globally toggle or role-restrict specific Magic Studio tools (e.g., locking Text-to-Image for specific student groups while allowing Magic Write for staff). Contractual Data Shield Enforcement: Under the district's Canva for Education framework (NDPA/FERPA/COPPA compliant), all user prompts, uploads, and creations are contractually excluded from public model training cycles automatically. Automated Input Content Moderation: Canva features automated pre-execution filters that instantly intercept and block offensive, copyrighted, or policy-violating prompts before they reach the backend AI engine.
Mitigation steps available to limit identified security risks (§219.22(b)(1))	Mandatory Configuration Change: Every user must manually navigate to Settings > Data Controls and turn off "Improve the model for everyone" to stop data training.

TAC §219.22(b) Requirement	Canva Response / Evidence
	Strict PII Ban: Complete institutional prohibition of inputting Personally Identifiable Information (PII), FERPA-protected student records, or sensitive financial data. Browser/Network Guardrails: Implementation of endpoint monitoring or web-filtering tools to block unapproved extensions interacting with the ChatGPT web interface.
Accuracy of system outputs — measurements of accuracy and relevance (§219.22(b)(2)(A))	No Empirical Accuracy Benchmarks: Canva does not provide objective factual accuracy metrics or statistical confidence intervals to institutional users for text or design layouts. Contextual / Aesthetic Validation: Because it functions as a design assistant rather than an analytical research repository, output success is measured contextually by educators based on stylistic alignment, visual layout clarity, and relevance to the curriculum. Hallucination Risk Threshold: Text generation tools (Magic Write) carry a low-to-moderate baseline risk of factual fabrication ("hallucinations"). Continuous human oversight and validation are strictly mandatory for text blocks prior to instructional use.
Operational performance metrics — model latency, uptime, and error rate (§219.22(b)(2)(B))	Latency: Text modifications and layout edits render within 1 to 3 seconds. Complex, compute-heavy tasks like Text-to-Image and Text-to-Video generation average 5 to 15 seconds depending on cloud queue density. Uptime: Tied directly to Canva's global infrastructure availability, supported by multi-region AWS environments. System status is tracked via Canva's public status dashboard, with no distinct custom SLAs provided for standard educational tenants. Error Rate: Heavily reliant on regional client network speeds and concurrent global API multi-tenant traffic. Processing hiccups are managed natively via front-end user notifications ("Something went wrong, please try again") and resolve gracefully without system crashes.
Description of the system's algorithms and how it makes decisions (§219.22(b)(3)(A))	Multi-Model Orchestration Engine: Merges proprietary computer vision algorithms (for layout generation, color palettes, and graphic alignment) with licensed foundational models (LLMs and diffusion networks via API cloud configurations). Decision-Making Mechanism: Operates mathematically through spatial proximity scoring for visual layouts and autoregressive vector prediction for language workflows. The system maps incoming text prompts to high-dimensional latent space representations to predict the

TAC §219.22(b) Requirement	Canva Response / Evidence
	highest-probability design arrangement, exercising no independent administrative logic or policy authority.
Data used to train the system's model (§219.22(b)(3)(B))	Training Corpus: Core models are pre-trained on expansive open-source web scraping datasets, text corpuses, public-domain media repositories, and Canva's own vast marketplace of licensed, commercial stock assets. Total Institutional Data Isolation: By virtue of the managed Canva for Education enterprise agreement, all district branding assets, classroom uploads, user prompts, and finished designs are completely sandboxed and omitted from Canva's AI retraining pipelines.
Availability of inputs and outputs for ongoing monitoring of decision-making (§219.22(b)(3)(C))	Centralized Admin Dashboard Visibility: The Technology Director / Central IT Admin possesses high-level structural oversight across the tenant workspace, including user activity reporting, sharing permission auditing, and cloud folder access tracking. Workspace Inspection and Audit Capabilities: Although administrators do not have a live streaming feed of raw prompt strings, they retain overarching administrative authority to review any project hosted within the school district's domain spaces, enabling rapid removal or locking of offending content.

Part C: Fairness & Bias Assessment

Question	Google Gemini: Response / Evidence
Has the vendor provided bias testing results for this system?	☒ Yes ☐ No ☐ In Progress Evidence: Google provides general model documentation and research updates regarding its ongoing alignment with Google's <i>AI Principles</i>. Furthermore, Google builds explicit Model Fairness Evaluation Toolkits directly into its enterprise developer ecosystems to allow organizations to test their own local data slices for bias.
Were any demographic disparities identified in accuracy or outcomes?	☒ Yes (describe below) ☐ No ☐ Unknown _____ Specifically regarding systemic over-correction. Google's Gemini models have faced notable public and academic scrutiny for algorithmic balancing errors. In an effort to mitigate historic, internet-wide racial and gender erasure, early tuning layers historically over-compensated, accidentally generating historically inaccurate or contextually absurd demographic outputs (e.g., generating racially diverse imagery for historic groups where it was factually incorrect).

Question	Google Gemini: Response / Evidence
What mitigation measures are in place for identified biases?	Context-Aware Filtering: Google continuously rolls out incremental foundational patches (e.g., transition to Imagen 3 / upgraded Gemini architectures) designed to balance demographic representation with precise contextual and historical accuracy. Double-Check Grounding: The "Double-Check Response" tool allows users to instantly execute a secondary live Google Search overlay, visually highlighting which statements are supported by external sources and which are unverified or potentially biased model hallucinations.
What student or staff subgroups are most affected by this system's outputs?	Most Negatively Affected: Students and staff working within highly specialized historical, theological, or hyper-localized cultural subjects where the model's generalized diversity-injection rules may distort factual accuracy. Most Positively Affected: heavy users of the Google Workspace environment (Docs, Slides, Sheets) who require assistive writing tools to bridge accessibility gaps, such as speech-to-text generation or real-time document summarization.
Has an independent equity audit been performed? (Required for Tier 3)	☐ Yes ☒ No ☐ Planned — Date: _____ While Google partners extensively with academic red-teams and civil-society groups to pressure-test their systems before major rollouts, they do <i>not</i> provide a standalone, independent institutional equity audit for your specific Google Workspace for Education domain deployment.

Question	Microsoft Copilot: Response / Evidence
Has the vendor provided bias testing results for this system?	☒ Yes ☐ No ☐ In Progress Evidence: Microsoft evaluates its generative models against its internal <i>Responsible AI Standard</i>. While Microsoft does not issue automated, real-time localized bias metrics directly to the consumer dashboard, they publish public Responsible AI Transparency Notes and Application Cards. These documents provide generalized evaluations of model fairness, representation limitations, and stereotyping vulnerabilities observed during red-teaming phases.
Were any demographic disparities identified in accuracy or outcomes?	☒ Yes (describe below) ☐ No ☐ Unknown _____ Structurally. Consistent with all large language models (LLMs) built on frontier transformer bases, there is an inherent performance disparity favoring Western demographics, Eurocentric cultural contexts, and standardized American English. The model performs with

Question	Microsoft Copilot: Response / Evidence
	higher accuracy and greater contextual nuance when processing prompts in English compared to regional dialects or non-English languages.
What mitigation measures are in place for identified biases?	Multi-layered System Guardrails: Microsoft implements foundational Reinforcement Learning from Human Feedback (RLHF) alongside active system-level metaprompts. These metaprompts inject invisible, structural safety guidelines into user queries before they hit the model, explicitly directing the system to reject the generation of hate speech, toxic stereotypes, or demographic exclusions. Grounding Architecture: Through Retrieval-Augmented Generation (RAG), Copilot binds its generation layer directly to your institution's own verified data repositories (SharePoint/OneDrive), minimizing the injection of broad internet-born biases.
What student or staff subgroups are most affected by this system's outputs?	Most Negatively Affected: Non-native English speakers, students/staff communicating in regional dialects, and individuals from non-Western cultural or historical backgrounds. The system may occasionally misinterpret nuance or penalize syntax variations. Most Positively Affected: Students with specific neurodivergent learning traits or physical disabilities who leverage Copilot's processing features to summarize overwhelming blocks of text or generate structural templates for writing assignments.
Has an independent equity audit been performed? (Required for Tier 3)	☐ Yes ☒ No ☐ Planned — Date: ____ Independent institutional equity audit is provided out-of-the-box by Microsoft. Microsoft performs external, third-party "red-team" security and fairness evaluations at the foundational model level (OpenAI base models), but the vendor does <i>not</i> execute an independent equity audit tailored to your school's specific tenant environment.

Question	ChatGPT: Response / Evidence
Has the vendor provided bias testing results for this system?	☒ Yes ☐ No ☐ In Progress Evidence: At the macro level. OpenAI publishes extensive system cards alongside major model releases (e.g., GPT-4o system cards). These reports explicitly outline the results of rigorous toxicity, stereotyping, and political bias tests executed prior to commercial launch. However, no ongoing bias testing metrics are customized or delivered to standard user accounts.

Question	ChatGPT: Response / Evidence
Were any demographic disparities identified in accuracy or outcomes?	☒ Yes (describe below) ☐ No ☐ Unknown _____ OpenAI openly documents that standard ChatGPT models carry structural biases toward Western viewpoints, standard English prose, and majority internet data footprints. Specific Educational Disparity: Research demonstrates an evaluation disparity where the model's tone-evaluation mechanisms may disproportionately flag or misjudge text written by English as a Second Language (ESL) students, misinterpreting their structural syntax variations as lower-quality or AI-generated work.
What mitigation measures are in place for identified biases?	Post-Training Moderation: OpenAI uses RLHF and specialized automated moderation endpoints to strip out flagrant demographic biases, hate speech, and explicit prejudices. User-Level Context Controls: Users have access to "Custom Instructions" to manually inject regional context or personal guardrails to steer the model away from defaulting to standard Western-centric assumptions.
What student or staff subgroups are most affected by this system's outputs?	Most Negatively Affected: ESL (English as a Second Language) students and students from historically marginalized socio-economic backgrounds whose community vernaculars are underrepresented in internet text corpuses. Most Positively Affected: General students utilizing the tool as an ungrounded, everyday writing brainstorm or coding assistant to lower baseline learning friction.
Has an independent equity audit been performed? (Required for Tier 3)	☐ Yes ☒ No ☐ Planned — Date: _____ Standard consumer ChatGPT accounts have not undergone an independent, education-specific institutional equity audit. OpenAI relies entirely on general internal testing and voluntary external research collaborations.

Question	Canva: Response / Evidence
Has the vendor provided bias testing results for this system?	☒ Yes ☐ No ☐ In Progress Evidence: Under the Canva Shield trust and safety umbrella, Canva documents that all its generative AI utilities (including Magic Media text-to-image and Magic Write text-to-text tools) undergo rigorous safety and debiasing evaluations before public deployment. Canva publicly shares its human-centric approach to testing visual fairness and has announced the development and open-sourcing of an internal "bias steering model" built to foster more inclusive representation. However, continuous or localized bias testing reports

Question	Canva: Response / Evidence
	are not delivered dynamically to individual school district tenants.
Were any demographic disparities identified in accuracy or outcomes?	☒ Yes (describe below) ☐ No ☐ Unknown_____ Canva acknowledges that foundational generative visual and text models carry historical internet bias, frequently defaulting to Eurocentric cultural traits, Western visual environments, and gender stereotypes for specific professions (e.g., defaulting to specific demographics when generating images of a "doctor" vs. a "teacher"). Specific Educational Disparity: For text-based workflows (Magic Write), the underlying language models rely on highly structured, standard English text corpuses. This creates an outcome disparity where creative writing or formatting outputs may feel less aligned with the conversational idioms, vernacular variations, or structural syntaxes typical of English as a Second Language (ESL) students.
What mitigation measures are in place for identified biases?	Canva utilizes an automated bias steering layer within its Text-to-Image pipeline. This engine acts as a backend prompt enhancer, dynamically injecting diverse demographic distributions (ethnicity, gender, age, ability) to ensure generative image sets look multi-cultural and balanced by default. Canva Shield Automated Input/Output Moderation: Automated algorithms scan incoming prompt text to intercept and block explicit discriminatory or hateful language. Simultaneously, output monitors intercept and filter out visual anomalies or biased edge cases before they render on the user's screen. Educational Community Feedback Loop: Canva features localized reporting mechanisms ("Report this image/text") allowing educators and students to flag culturally insensitive or stereotypical outputs, which feeds directly into the vendor's model refinement loop.
What student or staff subgroups are most affected by this system's outputs?	Most Negatively Affected: Students and staff attempting to generate hyper-specific, authentic visual materials for historically underrepresented or regional indigenous subcultures, as the generative media database may over-simplify or misrepresent localized traditions in favor of mainstream global media defaults. Most Positively Affected: Neurodivergent students, visual learners, and students facing fine-motor or spatial organization barriers. The multimodal nature of Canva's AI (e.g., text-to-image or instant slide

Question	Canva: Response / Evidence
	generation) allows these students to rapidly articulate complex academic concepts visually without being bottlenecked by manual layout or typing friction.
Has an independent equity audit been performed? (Required for Tier 3)	☐ Yes ☒ No ☐ Planned — Date: ____ While Canva for Education's data handling is thoroughly audited and certified by independent student safety organizations (such as iKeepSafe for strict FERPA/COPPA compliance and NDPA alignment), the individual generative algorithms within the Magic Studio suite have not undergone a third-party, education-specific institutional equity audit. Canva relies entirely on internal testing under Canva Shield guidelines and collaborative academic feedback.

Part D: Privacy & Data Security

Question	Google Gemini: Response / Evidence
What personally identifiable information (PII) does this system collect or process?	Processes Google Workspace institutional login profiles (Google Admin managed Entra/Workspace accounts), names, educational email addresses, IP endpoints, and internal organizational contents pulled contextually from Google Docs, Sheets, Slides, Gmail, or Drive files.
Is student data involved? If yes, is FERPA compliance documented?	☒ Yes — FERPA compliant ☐ No student data ☐ Under review
Is data used to train the vendor's AI models?	☐ Yes (describe consent mechanism) ☒ No ☐ Unknown
What cybersecurity controls does the vendor maintain? (Request NIST SP 800-53 documentation)	Google maintains enterprise-grade alignments mapped meticulously to NIST SP 800-53 parameters, verified structurally through FedRAMP Moderate certifications, SOC 2/3 reviews, and ISO/IEC 27001, 27017, and 27018 compliance certifications.
Data breach notification SLA in contract?	☒ Yes — Hours: ____ ☐ No ☐ Pending
Does the vendor allow student-level data export to fulfill FERPA/TEC 26.004(b) records requests, including AI-generated records?	☒ Yes ☐ No — issue to resolve ☐ Partial
Is Data Minimization in place? (TAC §219.11(h)) Only PII necessary for operations collected.	☒ Yes ☐ No — remediation needed

Question	Microsoft Copilot: Response / Evidence
What personally identifiable information (PII) does this system collect or process?	Processes institutional login credentials (Microsoft Entra ID), names, school email addresses, IP addresses, and any user-submitted PII or institutional data embedded

Question	Microsoft Copilot: Response / Evidence
	within prompts or attached files (such as OneDrive documents, emails, or SharePoint files).
Is student data involved? If yes, is FERPA compliance documented?	☒ Yes — FERPA compliant ☐ No student data ☐ Under review
Is data used to train the vendor's AI models?	☐ Yes (describe consent mechanism) ☒ No ☐ Unknown
What cybersecurity controls does the vendor maintain? (Request NIST SP 800-53 documentation)	Microsoft maintains a comprehensive framework mapped tightly to NIST SP 800-53 controls, verified independently via FedRAMP Moderate/High authorizations, SOC 1/2/3 compliance matrices, and ISO/IEC 27001/27018 certifications.
Data breach notification SLA in contract?	☒ Yes — Hours: ____ ☐ No ☐ Pending
Does the vendor allow student-level data export to fulfill FERPA/TEC 26.004(b) records requests, including AI-generated records?	☒ Yes ☐ No — issue to resolve ☐ Partial
Is Data Minimization in place? (TAC §219.11(h)) Only PII necessary for operations collected.	☒ Yes ☐ No — remediation needed

Question	ChatGPT: Response / Evidence
What personally identifiable information (PII) does this system collect or process?	Collects account names, personal or school email addresses, payment/billing details (if using a paid Plus tier), browser device cookies, IP addresses, and any unstructured PII explicitly written into prompts or uploaded within text/image files.
Is student data involved? If yes, is FERPA compliance documented?	☐ Yes — FERPA compliant ☐ No student data ☒ Under review No student data (<i>Crucial Compliance Warning:</i> Because this is a standard account, it does not carry an institutional FERPA agreement out of the box. Institutional policy must completely bar the entry of student-protected datasets).
Is data used to train the vendor's AI models?	☒ Yes (describe consent mechanism) ☐ No ☐ Unknown By default, standard consumer accounts consent to data training upon account creation. Consent Mechanism / Mitigation: Users must manually click through Settings Data Controls and slide off "Improve the model for everyone" to opt out of future model training loops.
What cybersecurity controls does the vendor maintain? (Request NIST SP 800-53 documentation)	OpenAI maintains SOC 2 Type II compliance and foundational data encryption parameters at rest (AES-256) and in transit (TLS 1.2+). However, NIST SP 800-53 controls are typically reserved for Enterprise/Team plans, meaning standard consumer account frameworks may require supplementary institutional documentation requests.

Question	ChatGPT: Response / Evidence
Data breach notification SLA in contract?	☐ Yes — Hours: ____ ☒ No ☐ Pending
Does the vendor allow student-level data export to fulfill FERPA/TEC 26.004(b) records requests, including AI-generated records?	☐ Yes ☐ No — issue to resolve ☒ Partial Individual users can manually download and export their specific account data footprint directly via their local dashboard, but central IT administrators lack any tenant-wide ability to force-export data.
Is Data Minimization in place? (TAC §219.11(h)) Only PII necessary for operations collected.	☐ Yes ☒ No — remediation needed Standard consumer models default to retaining and processing complete contextual string inputs for downstream model optimization cycles unless explicitly disabled manually.

Question	Canva: Response / Evidence
What personally identifiable information (PII) does this system collect or process?	Collects core directory details required for provisioning accounts under the managed district tenant, including user account names, school-issued email addresses, and structural group assignments (classes/roles). It also ingests technical telemetry data (browser/device profiles, IP addresses, authentication security logs) and any unstructured PII that staff or students explicitly insert into design text layouts, presentation prompts, graphic layers, or multimedia uploads.
Is student data involved? If yes, is FERPA compliance documented?	☒ Yes — FERPA compliant ☐ No student data ☐ Under review Evidence: Unlike standard consumer accounts, our managed Canva for Education framework is backed by a signed National Data Privacy Agreement (NDPA) or a localized Data Processing Addendum (DPA). Canva is an official signatory of the Student Data Privacy Consortium (SDPC) model, providing documented alignment with FERPA and COPPA safety requirements. Student data is strictly walled off within the district's private cloud workspace.
Is data used to train the vendor's AI models?	☐ Yes (describe consent mechanism) ☒ No ☐ Unknown Evidence / Mitigation: Under the protected <i>Canva Shield</i> protocols designated for verified K-12 institutional accounts, Canva contractually guarantees that no student or staff content, prompt inputs, branding assets, or completed multimedia designs are utilized to train their proprietary or third-party partner AI models. Model training exclusions are enforced automatically tenant-wide at the admin level, removing the need for individual manual user opt-outs.
What cybersecurity controls does the vendor maintain? (Request NIST SP 800-53 documentation)	Canva maintains institutional-grade security architecture verified by independent audits, including SOC 2 Type II and ISO/IEC 27001 certifications. All data payloads are encrypted utilizing industry-standard AES-256 protocols at rest and TLS 1.2+ protocols while in transit.

Question	Canva: Response / Evidence
	While Canva aligns its corporate infrastructure security framework with the high-level control families of NIST SP 800-53, specific mapping documentation for localized K-12 environments must be requested directly from the Canva Education Procurement Compliance division via our established account channel.
Data breach notification SLA in contract?	☒ Yes — Hours: ____ ☐ No ☐ Pending
Does the vendor allow student-level data export to fulfill FERPA/TEC 26.004(b) records requests, including AI-generated records?	☐ Yes ☐ No — issue to resolve ☒ Partial Individual users can manually download and export their specific account data footprint directly via their local dashboard, but central IT administrators lack any tenant-wide ability to force-export data.
Is Data Minimization in place? (TAC §219.11(h)) Only PII necessary for operations collected.	☒ Yes ☐ No — remediation needed

Part E: Human Oversight & Transparency

Question	Google Gemini: Response / Evidence
Can a human meaningfully review and override the AI system's output before action is taken?	☒ Yes ☐ No ☐ Partial — describe: ____
How will students, families, or employees be notified when this AI system influences a decision about them?	Notification is satisfied through the district's annual technology disclosure agreements and the Google Workspace for Education core privacy announcements. The AI acts only as an operational drafting aid and does not autonomously make enrollment, grading, or disciplinary decisions.
For public-facing AI: Does the system identify itself as AI before interacting with users? (HB 149 / BCC §552.051)	☐ Yes ☐ No — required ☒ N/A (not public-facing)
Is the district's public AI use policy updated to include this system?	☒ Yes ☐ No — Update required by: ____
Can the AI explain why it produced a given output (explainability)?	☐ Full ☒ Partial ☐ None
Is staff training in place to interpret and critically evaluate AI outputs?	☒ Yes ☐ No — Training by: ____
Is a Kill-Switch / emergency disable protocol documented for this system? (TAC §219.11(c))	☒ Yes — documented at: ____ ☐ No — must be established before deployment
Is ongoing accuracy monitoring in place throughout the AI lifecycle? (TAC §219.11(e), §219.22(b)(2))	☒ Yes — describe: ____ ☐ No — monitoring plan required

Question	Microsoft Copilot: Response / Evidence
Can a human meaningfully review and override the AI system's output before action is taken?	☒ Yes ☐ No ☐ Partial — describe: ____
How will students, families, or employees be notified when this AI system influences a decision about them?	The tool does not directly evaluate individuals or automate administrative decisions, notification occurs at the policy level. The district lists Copilot in its Acceptable Use Policy (AUP). If a staff member uses it to synthesize high-level parent communications or program materials, standard school communication headers or footers can include a basic transparency notice.
For public-facing AI: Does the system identify itself as AI before interacting with users? (HB 149 / BCC §552.051)	☐ Yes ☐ No — required ☒ N/A (not public-facing)
Is the district's public AI use policy updated to include this system?	☒ Yes ☐ No — Update required by: ____
Can the AI explain why it produced a given output (explainability)?	☐ Full ☒ Partial ☐ None
Is staff training in place to interpret and critically evaluate AI outputs?	☒ Yes ☐ No — Training by: ____
Is a Kill-Switch / emergency disable protocol documented for this system? (TAC §219.11(c))	☒ Yes — documented at: ____ ☐ No — must be established before deployment
Is ongoing accuracy monitoring in place throughout the AI lifecycle? (TAC §219.11(e), §219.22(b)(2))	☒ Yes — describe: ____ ☐ No — monitoring plan required

Question	ChatGPT: Response / Evidence
Can a human meaningfully review and override the AI system's output before action is taken?	☒ Yes ☐ No ☐ Partial — describe: ____
How will students, families, or employees be notified when this AI system influences a decision about them?	The district's public AI policy explicitly states that standard consumer-grade AI systems are classified as provisional brainstorming utilities. Staff are barred from inputting personal student dossiers to generate custom behavioral or administrative decisions, eliminating hidden algorithmic influence.
For public-facing AI: Does the system identify itself as AI before interacting with users? (HB 149 / BCC §552.051)	☐ Yes ☐ No — required ☒ N/A (not public-facing)
Is the district's public AI use policy updated to include this system?	☐ Yes ☒ No — Update required by: ____
Can the AI explain why it produced a given output (explainability)?	☐ Full ☐ Partial ☒ None
Is staff training in place to interpret and critically evaluate AI outputs?	☒ Yes ☐ No — Training by: ____
Is a Kill-Switch / emergency disable protocol documented for this system? (TAC §219.11(c))	☒ Yes — documented at: ____ ☐ No — must be established before deployment

Question	ChatGPT: Response / Evidence
Is ongoing accuracy monitoring in place throughout the AI lifecycle? (TAC §219.11(e), §219.22(b)(2))	☒ Yes — describe: ____ ☐ No — monitoring plan required

Question	Canva: Response / Evidence
Can a human meaningfully review and override the AI system's output before action is taken?	☒ Yes ☐ No ☐ Partial — describe: ____
How will students, families, or employees be notified when this AI system influences a decision about them?	The district's central AI policy explicitly classifies the Canva platform as an interactive digital design assistant and multimedia presentation utility. Because its AI functions are restricted exclusively to visual layouts, template structuring, and personal asset generation, the platform has zero intersection with, or capability to influence, any evaluative or administrative decision-making processes. Staff are structurally blocked from processing student evaluation metrics or behavioral portfolios through Canva, eliminating any hidden algorithmic influence over student outcomes.
For public-facing AI: Does the system identify itself as AI before interacting with users? (HB 149 / BCC §552.051)	☐ Yes ☐ No — required ☒ N/A (not public-facing)
Is the district's public AI use policy updated to include this system?	☐ Yes ☒ No — Update required by: ____
Can the AI explain why it produced a given output (explainability)?	☐ Full ☐ Partial ☒ None
Is staff training in place to interpret and critically evaluate AI outputs?	☒ Yes ☐ No — Training by: ____
Is a Kill-Switch / emergency disable protocol documented for this system? (TAC §219.11(c))	☒ Yes — documented at: ____ ☐ No — must be established before deployment
Is ongoing accuracy monitoring in place throughout the AI lifecycle? (TAC §219.11(e), §219.22(b)(2))	☒ Yes — describe: ____ ☐ No — monitoring plan required

Part F: HSAI Requirements Addendum (Tier 3 Only)

Complete this section for any AI system classified as Tier 3 (HSAI). This section directly tracks the mandatory requirements of TAC §§219.22 and 219.24.

Requirement	Legal Basis	Status
Written risk assessment completed before deployment or material change	TAC §219.22(a)-(b)	☐ Complete ☐ In Progress ☐ Not Started
Risk assessment documents the system's security risks, mitigation steps, performance metrics, and transparency information	TAC §219.22(b)	☐ Complete ☐ In Progress ☐ Not Started

Requirement	Legal Basis	Status
Designated employee (AI Risk Officer) reviewed the completed written risk assessment	TAC §219.22(c)(1)	☐ Complete ☐ Pending
AI Risk Officer approved or denied deployment and notified the executive head	TAC §219.22(c)(2)	☐ Approved — Date: ____ ☐ Denied ☐ Pending
Written risk assessment and all relevant documents retained per state records retention schedule	TAC §219.22(d)	☐ Yes — Location: ____ ☐ No
HSAI-specific Acceptable Use Policy (AUP) adopted identifying acceptable uses and limitations	TAC §219.24(b)	☐ Yes — AUP document: ____ ☐ No — Required before deployment
All employees using this HSAI system trained on the specific AUP	TAC §219.24(b)	☐ Yes ☐ No — Training by: ____
Risk-specific training provided to all employees who access, use, or manage this HSAI system	TAC §219.24(c)	☐ Yes ☐ No — Training by: ____
Vendor contractually required to implement the NIST AI Risk Management Framework	TAC §219.24(d)	☐ Yes ☐ No — Required for HSAI
Formal demographic accuracy audit completed by independent reviewer	Best practice / SB 1964	☐ Complete ☐ Scheduled: ____ ☐ Not started
Transparency documentation (model card) obtained from vendor	Best practice / SB 1964	☐ Complete ☐ In progress ☐ Not available
Executive head formally approved deployment	Local governance	☐ Approved — Date: ____ ☐ Pending
Individual notice mechanism tested and operational	SB 1964 / TAC §219.11(g)	☐ Tested ☐ Pending ☐ Not applicable
Grievance/redress process for affected individuals documented and communicated	TAC §219.11(f)	☐ Yes — See Section ____ ☐ No
Semi-annual review schedule established	Best practice / TAC §219.11(k)	☐ Yes — Dates: ____ ☐ No
Ongoing accuracy monitoring plan established for system lifecycle	TAC §219.11(e), §219.22(b)(2)	☐ Yes ☐ No — Required
AI impact assessment considered/conducted (encouraged for districts)	TAC §219.23(e)	☐ Completed ☐ In progress ☐ Chose to defer — reason: ____

Signature	Name / Date
Technology Director Signature / Date:	
Superintendent Signature / Date:	

SECTION 8: NOTICE & TRANSPARENCY REQUIREMENTS

8.1 Individual Notice

When a **Tier 2 or** Tier 3 AI system is used in a decision affecting a student's or employee's rights, benefits, privileges, or disciplinary status, the district shall provide notice that:

- Clearly identifies that an AI system was used or considered in the decision
- Describes in plain language what the AI system does and how it informed the decision
- States the individual's right to request human review of the decision
- Provides contact information to submit questions or a complaint

8.2 Public-Facing AI Disclosure (HB 149 / BCC §552.051)

Under HB 149 (Texas Responsible AI Governance Act), governmental entities that make an AI system available to interact with the public must disclose to each user, before or at the time of interaction, that they are interacting with an AI system regardless of whether it would be obvious to a reasonable person. All district chatbots, virtual assistants, and AI-powered communication tools on public-facing platforms (websites, parent portals, etc.) must include this disclosure before any interaction begins.

8.3 AI-Generated Records and Public Information Requests

AI-generated records may be subject to the Texas Public Information Act (Government Code Ch. 552). The district shall treat AI-generated records according to their applicable classification:

- Website chatbot conversation histories: classified as General Correspondence — may be subject to a public information request (PIR)
- AI meeting transcription outputs: classified as Staff Meeting Minutes — subject to PIR
- AI-generated lesson plans or professional development content: review with legal counsel for PIR applicability
- AI tutoring or intervention records: subject to FERPA and TEC §26.004(b) parent records requests

The district shall consult with legal counsel regarding records retention obligations for AI-generated content. The Technology Director / AI Risk Officer shall ensure vendors can produce AI-generated records in response to PIR and FERPA requests.

SECTION 9: GOVERNANCE & OVERSIGHT STRUCTURE

9.1 Roles & Responsibilities

Note: Regarding the AI Risk Officer: TAC §219.21(a) requires school districts, as local governmental entities, to designate an employee as the AI Risk Officer. This may be an existing employee, such as the Technology Director, with this responsibility added to their current duties. A separate hire is not required.

Role	AI Governance Responsibilities
Superintendent	Formally approve Tier 3 AI deployments; receive annual AI governance report; adopt and revise AI framework policy. Overall accountability for district AI compliance; Tier 3 deployments; designate AI Governance Lead; receive notification from the Technology Director of HSAI deployment approvals or denials (TAC §219.22(c)).
Technology Director (AI Governance Lead / AI Risk)	Designated as the district's AI Risk Officer (if assigned); maintain AI system inventory; lead risk assessments for HSAI systems; review completed risk

Role	AI Governance Responsibilities
Officer** Per TAC §219.21(a), a designated district employee must serve as AI Risk Officer. The Technology Director may fill this role.	assessments and approve or deny HSAI deployment; notify the executive head of HSAI decisions; coordinate vendor compliance; manage DIR reporting; lead annual framework review.
Campus Principals	Ensure campus-level AI use complies with this framework; report any unauthorized AI use to Technology Director / AI Risk Officer; support staff training completion.
Chief Financial Officer	Ensure AI procurement contracts include required compliance provisions; coordinate with legal counsel on vendor agreements.
Human Resources Director	Ensure AI systems used in HR decisions comply with heightened scrutiny requirements; manage staff AI use training programs.
Legal Counsel	Review heightened scrutiny AI contracts; advise on SB 1964, HB 3512, and TAC 219 compliance; manage complaints and appeals involving AI; advise on public records obligations for AI-generated content.
Teaching & Learning Staff	Evaluate AI instructional tools for pedagogical quality and equity; report concerns to campus principal or Technology Director.

9.2 Annual AI Governance Review

The Technology Director / AI Risk Officer shall lead an annual review of this framework each year. The review shall include:

- **Audit of the AI System Inventory for accuracy and completeness**
- **Risk assessment updates for all active Tier 3 (HSAI) systems**
- **Review of any complaints or incidents involving AI systems**
- **Assessment of staff AI literacy and training completion rates**
- **Review of DIR guidance and any updates to the statewide AI Code of Ethics (TAC §219.11(a))**
- **Verification that all HSAI risk assessment records are retained per state records retention schedule (TAC §219.22(d))**
- **Review of DIR Prohibited Technologies List and Governor's Executive Order for any newly prohibited AI tools**
- **Recommendation to the Board for any framework updates**

9.3 Redress for Consequential AI Decisions (TAC §219.11(f))

Pursuant to TAC §219.11(f), the district shall maintain a method of redress for individuals who believe an AI system has made a consequential decision that unfairly impacts them in a material way. The district may fulfill this obligation through existing grievance processes outlined in policies FNG(LOCAL), DGBA(LOCAL), and GF(LOCAL), provided those processes:

- Allow the affected individual to request human review of an AI-influenced decision
- Ensure a qualified staff member, not the AI system, makes the final determination on the grievance
- Document the AI system's role in the original decision and the outcome of the grievance
- Are communicated to affected individuals as part of the AI notice required under Section 8.2

Contact for AI Redress Requests: Christopher Escarsega, AI Risk Officer, 915-765-3036, Escarsegac@tisd.us

SECTION 10: STAFF & STUDENT AI USE GUIDELINES

10.1 Acceptable Use — All AI Tools

All users of district AI systems (staff, students, and contractors) must:

- Use AI tools only for authorized educational or administrative purposes
- Never input sensitive personal data (SSNs, medical records, discipline records) into AI tools not approved for such use
- Never input PII or confidential data into unauthorized or publicly available AI systems (TAC §219.11(h)(4)(C))
- Always apply critical judgment to AI-generated content; never accept AI outputs as final without review
- Report suspected AI errors, bias, or misuse to a supervisor or the Technology Director
- Not attempt to circumvent AI system controls, filters, or safety features

10.2 Staff AI Use — Additional Requirements

- Staff may not use AI systems to make unilateral decisions about students' academic standing, discipline, or special education eligibility
- AI-generated assessments, reports, or communications must be reviewed and approved by a qualified staff member before being acted upon or shared
- Staff must complete district AI literacy training before using any Tier 2 or Tier 3 system
- AI use in grading must be disclosed to students and families
- Staff must not use AI tools to discuss or provide content related to human sexuality, sexual orientation, or gender identity unless the tool has been vetted under TEC §28.004 and district policy

10.4 AI Literacy Training Requirements (HB 3512)

HB 3512 (eff. Sep 1, 2025) requires annual DIR-certified AI training for the cybersecurity officer (and any elected/appointed officials who have access to a local government computer system and utilize it to perform at least 25% of their duties) , in addition to existing cybersecurity training.

SECTION 11: OPERATIONAL AI SAFETY PROTOCOLS

11.1 Emergency Disable Protocol (TAC §219.11(c))

Pursuant to TAC §219.11(c)(2)(C), the district must ensure AI systems can be disabled until harmful or inaccurate outputs can be remedied. For each deployed AI system, the AI risk officer shall document:

- The specific mechanism or procedure for disabling the system (e.g., vendor contact, API key revocation, network block)
- The personnel authorized to initiate the disablement protocol
- The threshold conditions that trigger immediate disablement (e.g., biased outputs, safety violation, data breach)
- The notification chain when a system is disabled (executive head, affected campus principals)
- The review process before a disabled system may be reactivated

AI System Name	Disable Mechanism	Authorized Personnel	Reactivation Approval
All AI apps list on AI Framework.	Content Keeper	AI Risk Officer	AI Risk Officer

AI System Name	Disable Mechanism	Authorized Personnel	Reactivation Approval

11.2 AI Accuracy Monitoring Throughout Lifecycle (TAC §219.11(e), §219.22(b))

Pursuant to TAC §219.11(e) and §219.22(b), the district shall formalize processes for monitoring AI system accuracy before deployment and throughout the AI lifecycle. The AI Team shall establish a monitoring plan for each active AI system that includes:

- Pre-deployment accuracy baseline testing
- Periodic spot-checking of AI outputs
- A defined process for staff to report suspected AI errors or inaccurate outputs
- Escalation procedures when accuracy degradation is detected
- Documentation of monitoring results in the AI System Inventory

SECTION 12: BIOMETRIC DATA & STUDENT CONSENT REQUIREMENTS

This section establishes the district's requirements for AI systems that collect or process biometric data or sensitive student information, pursuant to TEC §26.009, SB 12, and HB 149.

12.1 Biometric Data — Prior Written Parental Consent Required

The district shall obtain prior written parental consent before any AI system collects, stores, or uses biometric data, including faceprints, voiceprints, or fingerprints, to identify students. (TEC §26.009(a)(2)(B); SB 12)

- The consent form shall clearly explain what biometric data is collected, how it is stored, who has access, and how it will be used
- AI tools using voice recognition for any purpose (including transcription that identifies individual speakers) shall be reviewed for biometric consent applicability
- No AI system shall collect student biometric data without documented prior written consent in the student's file
- The AI risk officer shall maintain a registry of all AI tools that process biometric data and confirm consent status

12.2 Health, Wellness, and Mental Health AI — Parental Consent Required

AI tools that function as health or wellness tools and capture psychological data used for evaluation or treatment of students require parental consent. (TEC §26.009(a)(1); SB 12)

- Before deploying any AI mental health tool, wellness survey platform, or social-emotional learning AI that captures psychological data, the district shall notify parents and obtain written consent
- These systems shall be classified as Tier 3 (HSAI) and subject to full heightened scrutiny protocol
- The district shall ensure these tools do not initiate or facilitate conversations about human sexuality, sexual orientation, or gender identity without compliance with [TEC §28.004](#)

12.3 Prohibited Biometric Identification Without Consent

The district shall not develop or deploy any AI system for the purpose of uniquely identifying a specific individual using biometric data, or for the targeted or untargeted gathering of images or other media from the Internet or other publicly available sources without the individual's consent, if such gathering would infringe on any right of the individual under the U.S. Constitution, the Texas Constitution, or state or federal law. ([HB 149](#))

SECTION 13: FRAMEWORK ADOPTION

Field	Entry
Board Meeting Date:	7/29/2026
Board Meeting Location:	19210 Cobb Ave, Tornillo, TX 79853
Agenda Item Number:	7.d.
Superintendent Signature:	
Technology Director Signature:	

Version History

Version	Date	Author	Summary of Changes
1.0			Initial adoption

APPENDIX A — AI SYSTEM CLASSIFICATION QUICK REFERENCE WORKSHEET

Use this worksheet to classify a proposed AI system before beginning the full risk assessment. Answer each question in sequence and stop at the first matching tier.

#	Classification Question (Gemini for Google)	If YES →	Answer
1	Does the AI system's output directly and automatically determine an outcome affecting a student or employee (e.g., automatic denial, automatic discipline trigger) without any human review step?	→ TIER 3 (HSAI)	☐ Yes ☒ No
2	Does the AI output serve as the "principal basis" for a human decision about an individual, even if a human technically makes the final call?	→ TIER 3 (HSAI)	☐ Yes ☒ No
3	Does the AI output inform a human decision-maker who has genuine discretion and authority to override the system, and this override capability is actually used in practice?	→ TIER 2 (Moderate)	☐ Yes ☒ No
4	Is the AI system purely assistive (helping with tasks like scheduling, drafting, or aggregating information) with no individual-level decision output?	→ TIER 1 (Low)	☒ Yes ☐ No
5	Does the AI system collect or process biometric data (faceprints, voiceprints, fingerprints) to identify students?	→ Biometric Consent Protocol required (Section 12) regardless of tier	☐ Yes ☒ No
6	Is this a health/wellness AI that captures psychological data for evaluation or treatment of students?	→ Tier 3 + Parental Consent required (Section 12.2)	☐ Yes ☒ No

#	Classification Question (Gemini for Google)	If YES →	Answer
7	Has the vendor been screened against the DIR Prohibited Technologies List ?	→ Must clear before any tier deployment	☒ Cleared ☐ Pending ☐ Prohibited

System Being Classified:	Final Tier Classification:
Classified By: Christopher Escarsega, AI Rick Officer	Date: 6/4/2026

#	Classification Question (Microsoft Copilot)	If YES →	Answer
1	Does the AI system's output directly and automatically determine an outcome affecting a student or employee (e.g., automatic denial, automatic discipline trigger) without any human review step?	→ TIER 3 (HSAI)	☐ Yes ☒ No
2	Does the AI output serve as the "principal basis" for a human decision about an individual, even if a human technically makes the final call?	→ TIER 3 (HSAI)	☐ Yes ☒ No
3	Does the AI output inform a human decision-maker who has genuine discretion and authority to override the system, and this override capability is actually used in practice?	→ TIER 2 (Moderate)	☐ Yes ☒ No
4	Is the AI system purely assistive (helping with tasks like scheduling, drafting, or aggregating information) with no individual-level decision output?	→ TIER 1 (Low)	☒ Yes ☐ No
5	Does the AI system collect or process biometric data (faceprints, voiceprints, fingerprints) to identify students?	→ Biometric Consent Protocol required (Section 12) regardless of tier	☐ Yes ☒ No
6	Is this a health/wellness AI that captures psychological data for evaluation or treatment of students?	→ Tier 3 + Parental Consent required (Section 12.2)	☐ Yes ☒ No
7	Has the vendor been screened against the DIR Prohibited Technologies List ?	→ Must clear before any tier deployment	☒ Cleared ☐ Pending ☐ Prohibited

System Being Classified:	Final Tier Classification:
Classified By: Christopher Escarsega, AI Rick Officer	Date: 6/4/2026

#	Classification Question (ChatGPT)	If YES →	Answer
1	Does the AI system's output directly and automatically determine an outcome affecting a student or employee (e.g., automatic denial, automatic discipline trigger) without any human review step?	→ TIER 3 (HSAI)	☐ Yes ☒ No
2	Does the AI output serve as the "principal basis" for a human decision about an individual, even if a human technically makes the final call?	→ TIER 3 (HSAI)	☐ Yes ☒ No

#	Classification Question (ChatGPT)	If YES →	Answer
3	Does the AI output inform a human decision-maker who has genuine discretion and authority to override the system, and this override capability is actually used in practice?	→ TIER 2 (Moderate)	☐ Yes ☒ No
4	Is the AI system purely assistive (helping with tasks like scheduling, drafting, or aggregating information) with no individual-level decision output?	→ TIER 1 (Low)	☒ Yes ☐ No
5	Does the AI system collect or process biometric data (faceprints, voiceprints, fingerprints) to identify students?	→ Biometric Consent Protocol required (Section 12) regardless of tier	☐ Yes ☒ No
6	Is this a health/wellness AI that captures psychological data for evaluation or treatment of students?	→ Tier 3 + Parental Consent required (Section 12.2)	☐ Yes ☒ No
7	Has the vendor been screened against the DIR Prohibited Technologies List ?	→ Must clear before any tier deployment	☒ Cleared ☐ Pending ☐ Prohibited

System Being Classified:	Final Tier Classification:
Classified By: Christopher Escarsega, AI Risk Officer	Date: 6/4/2026

APPENDIX B — SB 1964 / TAC 219 / HB 3512 COMPLIANCE CHECKLIST

For use by: (Technology Director / AI Risk Officer) | Items marked (*) are expressly required of school districts under TAC 219 or SB 1964.

Requirement	Legal Basis	Status
* AI Code of Ethics formally adopted	TAC §219.11(a)	☐ Complete ☒ In Progress ☐ Not Started
* Employee designated as AI Risk Officer (may be Technology Director with added duties)	TAC §219.21(a)	☒ Complete ☐ In Progress ☐ Not Started
* Process established to identify and inventory all HSAI systems	TAC §219.21(b)	☒ Complete ☐ In Progress ☐ Not Started
* Written risk assessment completed before HSAI deployment or material change	TAC §219.22(a)-(b)	☒ Complete ☐ In Progress ☐ Not Started
* Risk assessment documents security risks, performance metrics, and transparency info	TAC §219.22(b)	☒ Complete ☐ In Progress ☐ Not Started
* AI Risk Officer reviewed risk assessment and approved/denied HSAI deployment	TAC §219.22(c)(1)-(2)	☒ Complete ☐ In Progress ☐ Not Started
* Executive head notified of AI Risk Officer's HSAI deployment decision	TAC §219.22(c)(2)	☐ Complete ☒ In Progress ☐ Not Started
* Written risk assessment retained per state records retention schedule	TAC §219.22(d)	☒ Complete ☐ In Progress ☐ Not Started

Requirement	Legal Basis	Status
* HSAI-specific Acceptable Use Policy (AUP) adopted for each HSAI system	TAC §219.24(b)	☒ Complete ☐ In Progress ☐ Not Started
* All employees using each HSAI system trained on its specific AUP	TAC §219.24(b)	☐ Complete ☒ In Progress ☐ Not Started
* Risk-specific training provided to employees who access, use, or manage each HSAI	TAC §219.24(c)	☐ Complete ☒ In Progress ☐ Not Started
* Vendor contracts for HSAI systems require NIST AI RMF implementation	TAC §219.24(d)	☒ Complete ☐ In Progress ☐ Not Started
* DIR procurement requirements for HSAI systems followed	TAC §219.23(e)(1)	☒ Complete ☐ In Progress ☐ Not Started
AI impact assessment considered/conducted for HSAI deployment (encouraged for districts)	TAC §219.23(e)(2)-(3)	☐ Completed ☐ Deferred — reason:
* Individual notice mechanism in place for heightened scrutiny systems	SB 1964	☐ Complete ☐ In Progress ☒ Not Started
* Redress process for AI-affected individuals documented	TAC §219.11(f)	☒ Complete ☐ In Progress ☐ Not Started
* Disablement protocols documented for each active AI system	TAC §219.11(c)	☐ Complete ☒ In Progress ☐ Not Started
* Ongoing AI accuracy monitoring plan in place	TAC §219.11(e), §219.22(b)(2)	☐ Complete ☒ In Progress ☐ Not Started
* Data minimization policy in place	TAC §219.11(h)	☒ Complete ☐ In Progress ☐ Not Started
* PII prohibition in unauthorized AI systems communicated and enforced	TAC §219.11(h)(4)(C)	☒ Complete ☐ In Progress ☐ Not Started
* AI governance aligned with NIST AI Risk Management Framework	SB 1964	☒ Complete ☐ In Progress ☐ Not Started
* Vendor contracts include SB 1964 compliance provisions	SB 1964	☒ Complete ☐ In Progress ☐ Not Started
* Annual framework review scheduled	SB 1964 / TAC 219	☐ Complete ☐ In Progress ☐ Not Started
* Board of Trustees approved Tier 3 deployments	Local governance	☐ Complete ☐ In Progress ☐ Not Started
* DIR-certified AI training in place for all staff using computers ≥25% of duties (optional)	HB 3512; TEC §11.175	☐ Complete ☐ In Progress ☐ Not Started
* Cybersecurity coordinator completed DIR-certified cybersecurity AND AI training annually (and any elected/appointed officials who have access to a local government computer system and utilize it to perform at least 25% of their duties)	TEC §11.175(h-1); HB 3512	☐ Complete ☒ In Progress ☐ Not Started
* All public-facing AI systems disclose AI interaction before use	HB 149 / BCC §552.051	☐ Complete ☐ In Progress ☒ Not Started
* Prior written parental consent obtained for any student biometric data collection	TEC §26.009; SB 12	☐ Complete ☐ In Progress ☒ Not Started

Requirement	Legal Basis	Status
* Parental consent protocol in place for health/wellness AI capturing psychological student data	TEC §26.009; SB 12	☐ Complete ☐ In Progress ☒ Not Started
* All AI tools screened against DIR Prohibited Technologies List	DIR / Gov. Executive Order	☒ Complete ☐ In Progress ☐ Not Started
* AI vendor model disclosures reviewed for DIR Prohibited Technologies List compliance	Gov. Executive Order	☒ Complete ☐ In Progress ☐ Not Started
* Vendor contracts require student-level data export for FERPA/TEC §26.004(b) requests	FERPA; TEC §26.004(b)	☒ Complete ☐ In Progress ☐ Not Started
* AI-generated records retention policy reviewed with legal counsel	Texas Public Information Act	☐ Complete ☐ In Progress ☒ Not Started

Disclaimer & Adoption Record

Not Legal Advice This checklist is a compliance and operational reference compiled by Education Service Center Region 19 (ESC 19) for the School Systems it serves. It is not legal advice and does not establish an attorney-client relationship. Review with qualified legal counsel before adoption or any policy change. Statutory citations and effective dates were current as of the publication date; verify against the latest Texas Register, DIR publications, and TEA correspondence before relying on any provision. Local board policies and individual circumstances may require additional or different measures.