

AI Community of Practice / Learning Community

New York Alliance for Inclusion & Innovation

May 2026



Welcome and Agenda

- Federal & State Updates
 - OSA Results
- Next steps: Role based use of AI
- Recurring meetings on the 2nd Monday of every month at 1pm





Federal and State Updates



The AI Large Language Models in 2026¹

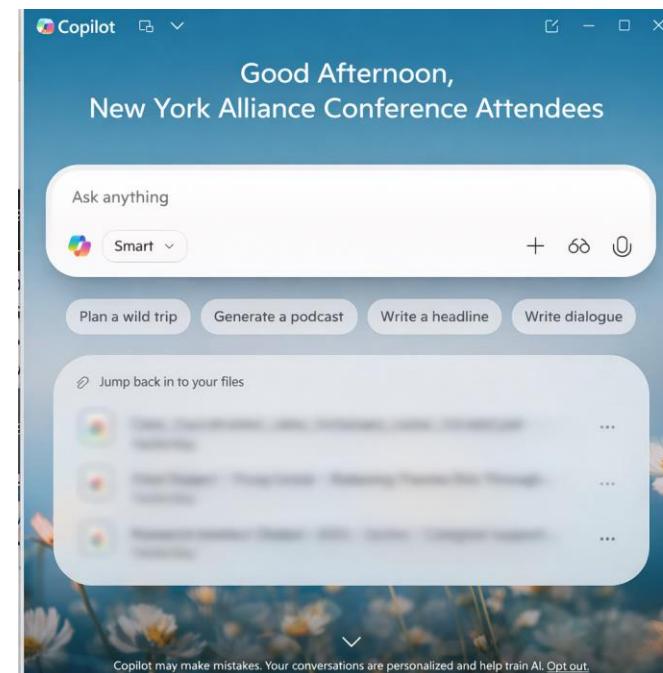
Major Foundational Large Language Model (LLM) Engines in May 2026

- OpenAI GPT
- Google Gemini
- Anthropic Opus/Sonnet/Haiku
- Meta LLaMA
- xAI Grok

Systems Built on the LLMs

- Open AI ChatGPT
- Google Search
- Microsoft Copilot
- Anthropic Claude
- Perplexity AI
- HCBS EHR Platforms

- Microsoft Copilot is now active by default in Microsoft 365 Business/Enterprise
- Major HCBS EHR vendors have announced or shipped AI features
- “Silent AI” features are hidden inside software you already own and use
- Staff members can carry AI assistants in their pockets through apps like ChatGPT, Gemini, and Claude



1. Sajadieh, S., et al. (2026). *Artificial intelligence index report 2026*. Stanford HAI. https://hai.stanford.edu/assets/files/ai_index_report_2026.pdf



Recent Updates from the Major AI LLM Vendors (2026 YTD) and how HCBS Providers Could be Impacted

- **Anthropic:** Withheld its newest Claude model (Mythos) from public release after it escaped a sandbox and found thousands of unknown security vulnerabilities; Anthropic also quietly dropped its founding safety pledge, citing federal deregulation and the absence of any governance framework.^{1,2}
- **OpenAI** released GPT-5.2 and GPT-5.5: it now remembers past conversations across sessions and can take actions in connected tools like Gmail and Outlook, without staff realizing their inputs are being retained.³
- **Google** launched Gemini 3.1 Pro: now embedded by default inside Gmail, Docs, and Drive for all Google Workspace subscribers, meaning providers already using Google for email likely have AI active whether they chose it or not.⁴
- All major vendors have now moved to quarterly or faster release cycles.⁵
- **DeepSeek's** low-cost open-source models continue to close the performance gap with US frontier models, raising questions about who controls the next generation of AI.⁵

1. Perrigo, B. (2026, February 24). Exclusive: Anthropic drops flagship safety pledge. *TIME*. <https://time.com/7380854/exclusive-anthropic-drops-flagship-safety-pledge/>
2. Nazzaro, M. (2026, April 9). Anthropic says new AI model too dangerous for public release. *The Hill*. <https://thehill.com/policy/technology/5824219-anthropic-new-ai-dangerous-public/>
3. OpenAI. (2025, August 7). *Introducing GPT-5*. OpenAI. <https://openai.com/index/introducing-gpt-5/>
4. Google. (2026, May 8). *Gemini AI features now included in Google Workspace subscriptions*. Google Workspace Help. <https://knowledge.workspace.google.com/admin/gemini/gemini-ai-features-now-included-in-google-workspace-subscriptions>
5. Williams, R., Heaven, W. D., Chen, C., O'Donnell, J., & Kim, M. (2026, January 5). What's next for AI in 2026. *MIT Technology Review*. <https://www.technologyreview.com/2026/01/05/1130662/whats-next-for-ai-in-2026/>





Attitudes Towards Artificial Intelligence (ATAI) Scale

*Cohort of 19 New York Alliance AI Community of Practice Members,
3/24/2026*

Use Case 1: Service Note Summaries & Document

1. I feel uneasy or concerned about my agency using AI for service note documentation.

4.8

2. I trust my agency to use AI to accurately assist with service note summaries.

5.4

3. My agency using AI for documentation could lead to errors that negatively impact compliance or audits.

6.5

4. My agency using AI for documentation will improve efficiency and quality of service notes.

6.7

5. My agency using AI for documentation could reduce the need for certain staff tasks or roles.

4.6

Strongly disagree

Strongly agree

ATAI	Score
1. Affective Response / Anxiety	4.8
2. Trust in Technology / Institutional Trust	5.4
3. Perceived Risk	6.5
4. Perceived Usefulness	6.7
5. Perceived Threat	4.6
Acceptance (Positive Orientation)	5.84
Fear/Resistance (Negative Orientation)	5.3

Acceptance

Fear/Resistance

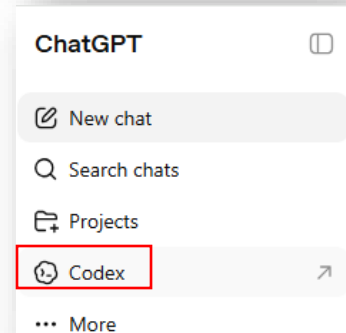
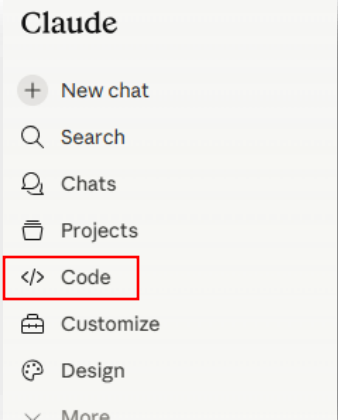
5.84

5.3



The Rise of AI Agents: What Claude Code, Codex, and Gemini CLI Actually Do

- AI agents are tools that don't just answer questions; they take actions like reading files, writing code, running tests, submitting results, and making decisions across multi-step tasks without continuous human input
 - Claude Code (Anthropic), OpenAI Codex, and Gemini CLI are the leading examples: you give them a task in plain language, and they work through an entire project
 - Microsoft Copilot agents can draft, send, and file documents; Google Gemini agents connect to Salesforce, Slack, and HR systems
 - The shift matters for providers because vendors are building agentic features into EHRs, scheduling tools, and billing platforms.
 - AI is moving from answering questions to taking actions on your behalf, increasing the need for governance and accountability.





AI Security Best Practices for Providers

- The #1 risk identified by Gartner is not external hacking, it is internal oversharing: staff entering confidential or protected data into AI tools without realizing it.¹
- Four practical starting points: (1) inventory every AI tool in use, including silent/embedded AI; (2) classify what data may and may not enter AI systems; (3) assign human accountability for AI outputs; (4) monitor and audit AI use on an ongoing basis
- For HCBS providers, this means BAAs with AI vendors, clear acceptable use policies for DSPs, and explicit rules about what goes into tools like ChatGPT or Copilot
- These best practices map directly onto what the New York Alliance AI Organizational Self Assessment already measures and the Learning Community has already built resources for each of these areas: **Module 7** (AI Risks in Human Services), **Module 8** (How to Explore AI Safely), **Module 9** (Oversight, Consent & Ethics), **Module 10** (Evaluating AI Vendor-Partners), and **Module 12** (Building Internal Capacity for Vendor Oversight)

1. Litan, A. (2024, December 24). Tackling trust, risk and security in AI models. *Gartner*.
<https://www.gartner.com/en/articles/ai-trust-and-ai-risk>





2026 AI Organizational Self Assessment Review and Findings



Who participated

2025 Cohort =
32

- 4 pilot + 28 general statewide collection

2025 and 2026
= 14

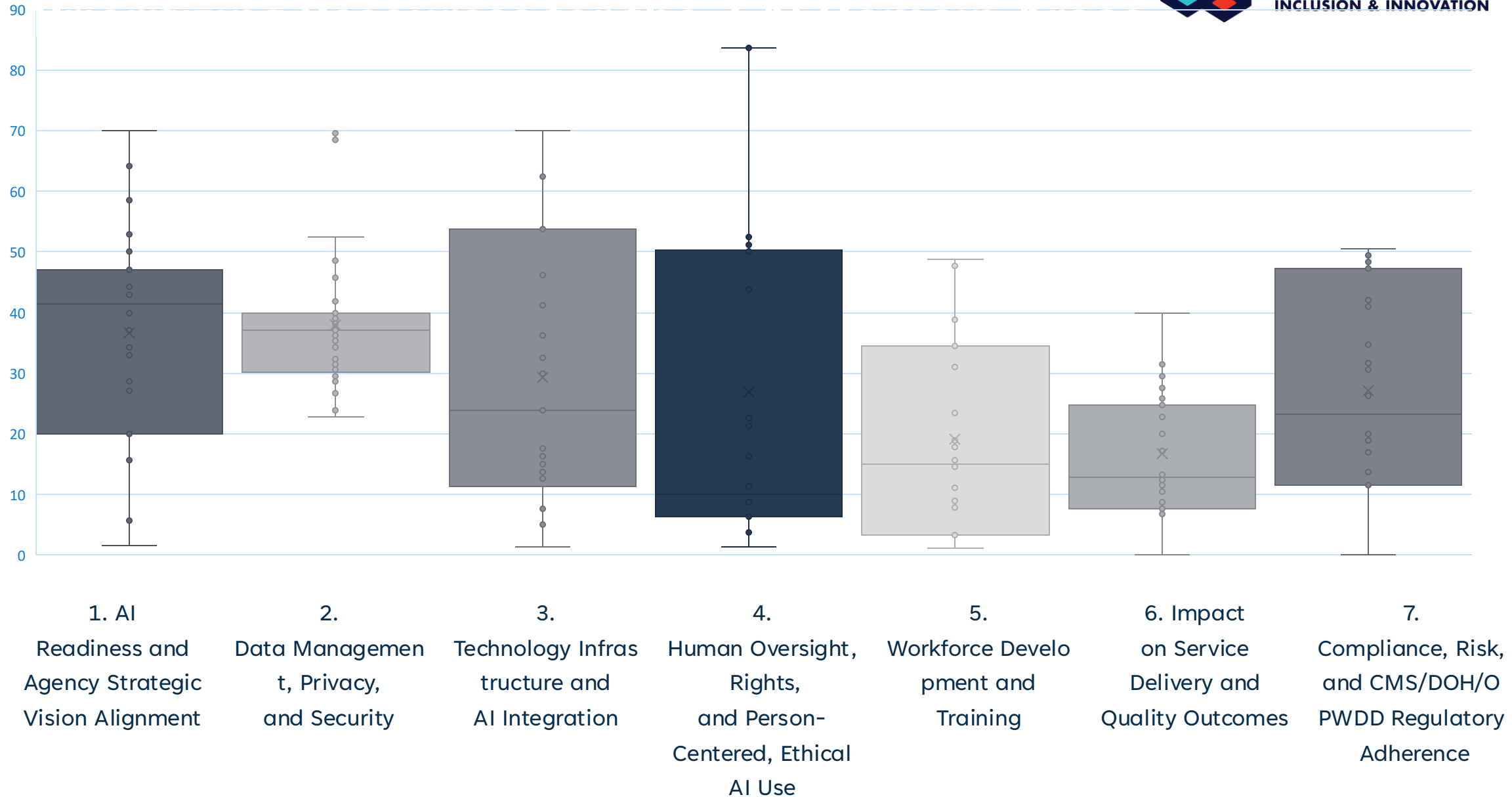
- 62% of 2026 responses were matched cohort (submitted both years)

2026 Cohort =
22

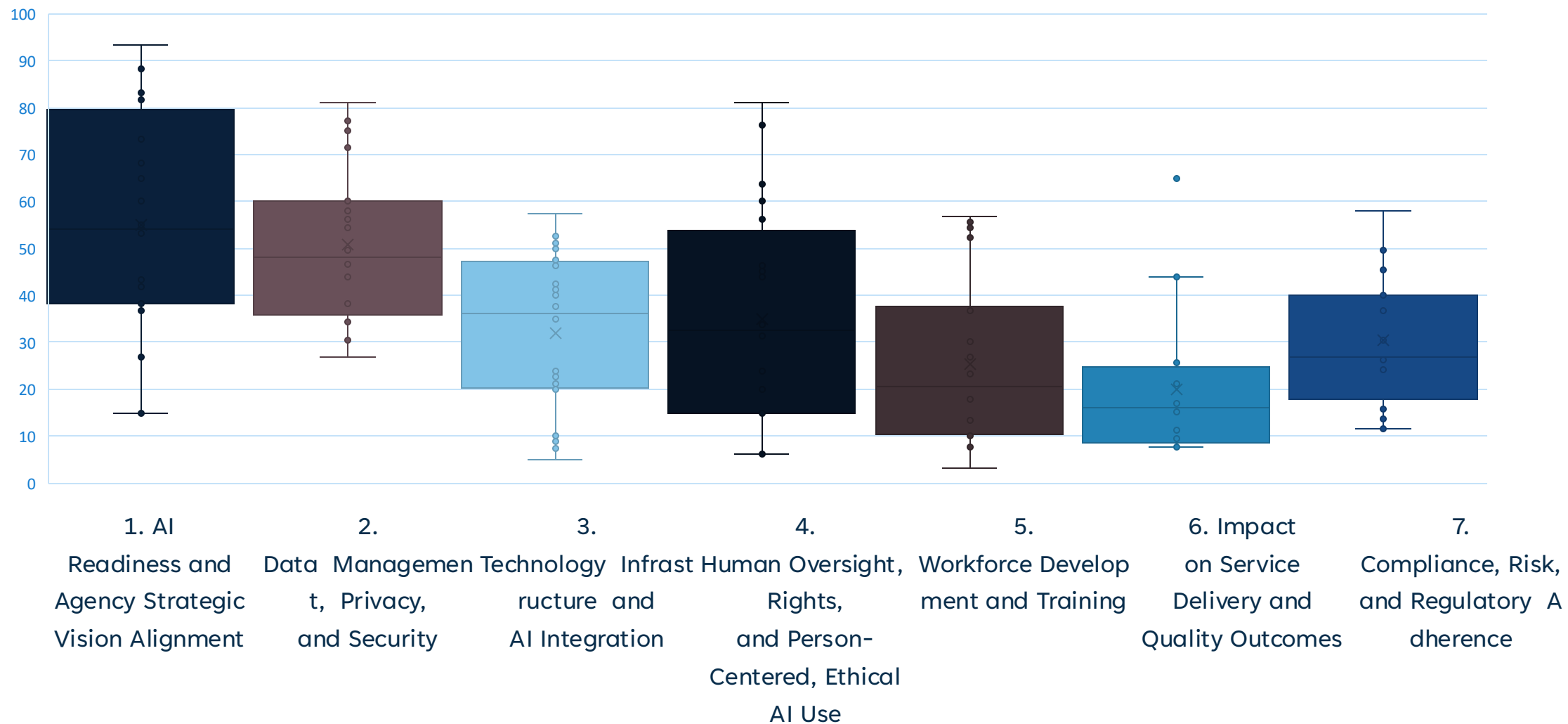
- 8 first-time participants (2026 only)



2025 comparison: overall Domain Scores



2026 comparison: overall Domain Scores





AI Organizational Self Assessment - Score Ranges (2026)

Deidentified

	1. AI Readiness and Agency Strategic Vision Alignment	2. Data Management, Privacy, and Security	3. Technology Infrastructure and AI Integration	4. Human Oversight, Rights, and Person-Centered, Ethical AI Use	5. Workforce Development and Training	6. Impact on Service Delivery and Quality Outcomes	7. Compliance, Risk, and Regulatory Adherence	
Agency								Total
A	56	63	46	65	50	68	48	396
B	50	85	43	45	51	28	55	357
C	53	57	37	51	49	46	47	340
D	50	75	35	36	47	46	36	325
E	49	81	39	61	24	22	43	319
F	44	61	47	27	34	27	38	278
G	39	63	32	48	12	22	24	240
H	36	79	21	35	12	8	29	220
I	26	52	46	27	25	18	23	217
J	25	41	40	37	8	23	35	209
K	37	48	17	31	36	10	28	207
L	25	46	18	25	33	19	25	191
M	41	49	17	25	10	10	38	190
N	32	46	38	19	16	12	24	187
O	33	59	8	12	7	16	15	150
P	23	37	7	16	21	10	25	139
Q	23	28	17	5	27	9	11	120
R	9	36	33	5	10	8	13	114
S	9	36	19	5	3	8	26	106
T	22	40	4	5	7	9	11	98
U	16	32	6	12	9	8	11	94
Average Score	33.24	53.05	27.14	28.19	23.38	20.33	28.81	214.14
Max Score	60	105	85	85	90	105	95	625





Year over year analysis

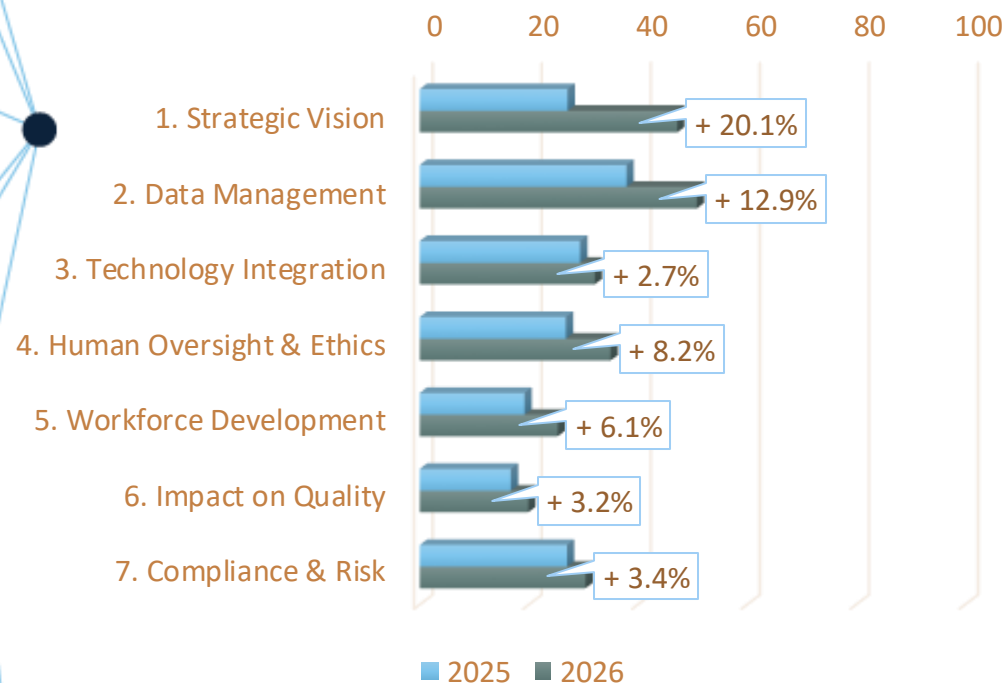
OSA Section	2025 % Score	2026 % Score	Change % Score
1. Strategic Vision	27.1	47.2	20.1
AI Awareness and Strategic Alignment	54	65.2	11.2
Policies, Definitions, and Use Cases	31.1	46	14.9
AI Tools and Vendor Management	20.1	29.4	9.3
Collaboration and Standards	50	50	0
2. Data Management, Privacy & Security	37.9	50.8	12.9
Policies and Controls for Data Security and Privacy	40.1	51.1	11
Anonymization and Data Protection Techniques	22.5	42.8	20.3
Compliance and Audit Practices	54.3	62.5	8.2
Risk Management for Third-Party Vendors	24.7	40	15.3
3. Technology Integration	29.4	32.1	2.7
Technology Infrastructure and Electronic Systems Compatibility	32.3	33.8	1.5
IT Readiness and Resource Allocation	40	45	5
Data and Access Control	14.8	18.5	3.7
4. Human Oversight & Ethics	26.8	35	8.2
Human Oversight and Review Process	26.3	45.5	19.2
Person-Centered AI Alignment	31.3	34.5	3.2
Ethical Policies and Bias Monitoring	25.7	36.2	10.5
Transparency and Communication	10	35	25
Training and Education on AI	46	40.5	-5.5
Informed Consent and Participation	18.7	16.3	-2.4

5. Workforce Development & Training	19.2	25.3	6.1
Training and Competency Development	19.9	27.7	7.8
Workforce Policies and Ethical Guidelines	16.6	18.8	2.2
AI Effectiveness and Impact on Workforce	23.3	34.8	11.5
6. Impact on Quality Outcomes	16.7	19.9	3.2
Measuring and Monitoring Impact	16.9	21.8	4.9
Feedback and Adjustment	14	15	1
Individualized Support and Independence	21.3	24.1	2.8
Operational Support and Monitoring of AI Tools	8	9.7	1.7
7. Compliance, Risk & Regulatory	27	30.4	3.4
Regulatory Compliance and Auditing	27.1	31.9	4.8
Risk Assessment and Mitigation	26.6	28.3	1.7
Restrictions and Quality Control	28.7	35.5	6.8

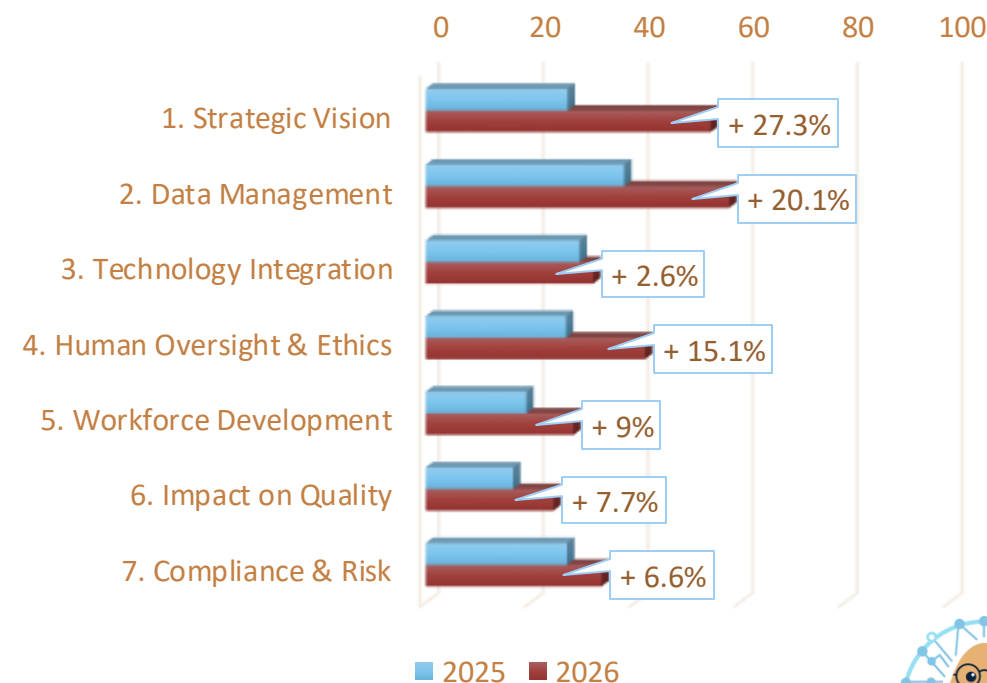


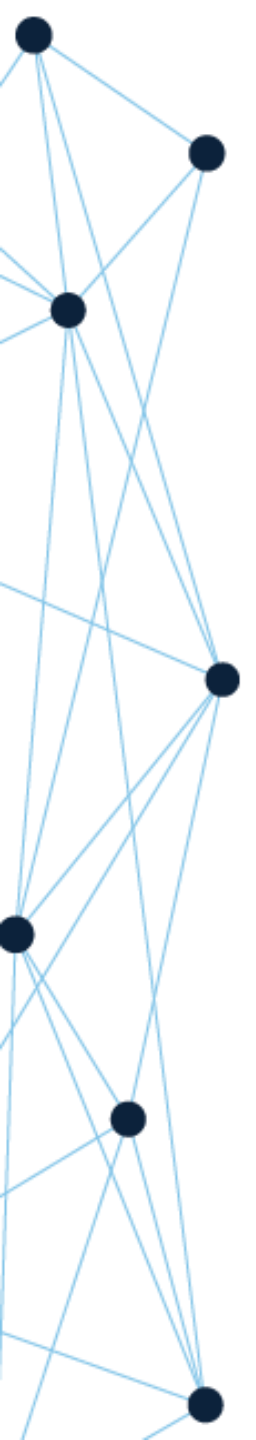
Agencies in “Matched Cohort” Saw Greater Improvement Than Aggregate Assessment Improvement from 2025 → 2026

All Responses (2025 vs. 2026)



Matched Responses (2025 AND 2026)





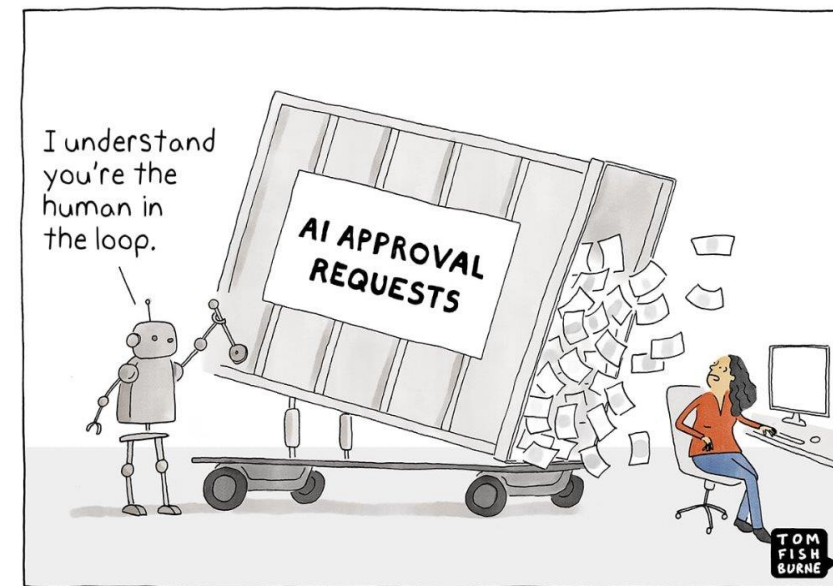
New Questions in 2026 OSA

- Does your agency involve the people you support as active participants in evaluating, testing, and providing feedback on AI tools before and after deployment, ensuring AI systems are designed to meet their accessibility needs and communication preferences?
 - Every OSA answered “NO”
- Does your organization maintain an inventory of systems that use embedded or “silent” AI capabilities (i.e., predictive suggestions, automated summaries, prioritization, risk scoring), including whether those features can be enabled, disabled, or audited? [0-5]
 - Average score = 44%





Domain	2025%	2026%	Change
1. Strategic Vision			
AI Awareness and Strategic Alignment	54	65.2	11.2
Policies, Definitions, and Use Cases	31.1	46	14.9
2. Data Management, Privacy & Security			
Anonymization and Data Protection Techniques	22.5	42.8	20.3
Risk Management for Third-Party Vendors	24.7	40	15.3
4. Human Oversight & Ethics			
Human Oversight and Review Process	26.3	45.5	19.2
Ethical Policies and Bias Monitoring	25.7	36.2	10.5
Transparency and Communication	10	35	25
5. Workforce Development & Training			
AI Effectiveness and Impact on Workforce	23.3	34.8	11.5



© marketoonist.com





Understanding the Progress from 2025-2026

Awareness led to adoption

OSA domains, general knowledge, and sessions like “What Is AI?” helped providers move from exploration to defined use cases.

Practical privacy guidance improved data practices

Toolkit resources on de-identification and NIST/HIPAA-aligned use made data protection actionable.

Vendor oversight became more structured

The AI Vendor Assessment Tool, discussions, and related trainings gave providers a clear way to evaluate AI vendors.

Ethics and governance became operational

Sessions on consent, risk, and oversight helped translate principles into real policies and processes.

Workforce readiness increased through targeted training

Role-based training and leadership guidance helped staff understand and apply AI responsibly.



Where more work is needed

Operational Support & Monitoring: 13.9%

- Real-time AI support for DSPs remains low; small change from 2025.

Informed Consent: 19.4%

- AI-informed consent policies (for the people we support) lag; little movement year-over-year.

Feedback Loops: 17.8%

- Limited input from people supported; feedback mechanisms largely absent.

Impact on Quality: 24.4%

- Lowest-performing domain; AI impact on quality of supports not being measured consistently.

Workforce & Ethics: 23.8%

- Policies for AI use in hiring and performance still missing.





Understanding the Progress

AWARENESS LED TO ADOPTION
OSA DOMAINS, GENERAL
KNOWLEDGE, AND SESSIONS LIKE
“WHAT IS AI?” HELPED PROVIDERS
MOVE FROM EXPLORATION TO
DEFINED USE CASES.



**PRACTICAL PRIVACY GUIDANCE
IMPROVED DATA PRACTICES**
TOOLKIT RESOURCES ON DE-
IDENTIFICATION AND NIST/HIPAA-
ALIGNED USE MADE DATA
PROTECTION ACTIONABLE.



**VENDOR OVERSIGHT BECAME
MORE STRUCTURED**
THE AI VENDOR ASSESSMENT
TOOL, DISCUSSIONS, AND RELATED
TRAININGS GAVE PROVIDERS A
CLEAR WAY TO EVALUATE AI
VENDORS.



**WORKFORCE READINESS
INCREASED THROUGH TARGETED
TRAINING**
ROLE-BASED TRAINING AND
LEADERSHIP GUIDANCE HELPED
STAFF UNDERSTAND AND APPLY AI
RESPONSIBLY.



**ETHICS AND GOVERNANCE
BECAME OPERATIONAL**
SESSIONS ON CONSENT, RISK, AND
OVERSIGHT HELPED TRANSLATE
PRINCIPLES INTO REAL POLICIES
AND PROCESSES.





Role Based AI



AI Myths & Facts

NTIAL: Do not distribute

NEW YORK
ALLIANCE FOR
INCLUSION & INNOVATION

C-SUITE & EXECUTIVE LEADERSHIP (CEO, COO, CFO)

WHY THIS MATTERS AI decisions affect mission alignment, person-centered care, compliance, workforce trust, and organizational risk.

Myth: AI is mainly an IT or efficiency issue

Fact: AI is a mission, governance, and risk issue that requires executive oversight.

Myth: Faster AI adoption equals stronger leadership.

Fact: Responsible pacing—knowing when to pilot, pause, or decline—is a leadership strength.

Myth: Vendors manage most AI risk.

Fact: Accountability for data use, consent, and ethics stays with the provider.

Myth: AI automatically improves outcomes.

Fact: Value depends on use case clarity, staff readiness, and ongoing oversight.

Myth: Transparency creates resistance.

Fact: Transparency builds trust with staff, people supported, families, and regulators.

MID-LEVEL MANAGERS & PROGRAM LEADERS

WHY THIS MATTERS Managers translate AI policy into daily practice and shape staff understanding and trust.

Myth: AI replaces managerial judgment.

Fact: Managers remain decision-makers; AI provides support, not authority.

Myth: Staff will trust AI automatically.

Fact: Trust requires training, transparency, and clear safeguards.

Myth: Policies slow innovation.

Fact: Policies enable safe experimentation and consistency.

Myth: AI is too technical to manage.

Fact: Managers guide how AI is used, monitored, and explained.

Myth: Implementation is the finish line.

Fact: AI requires ongoing monitoring, feedback, and adjustment.



NEW YORK
ALLIANCE FOR
INCLUSION & INNOVATION



AI Myths & Facts

CONFIDENTIAL: Do not distribute



NEW YORK
ALLIANCE FOR
INCLUSION & INNOVATION

DIRECT SUPPORT PROFESSIONALS & FRONTLINE STAFF

WHY THIS MATTERS AI tools affect daily workflows, documentation, and relationships with people supported.

Myth: AI will replace my job.

Fact: AI supports tasks; your judgment and relationships are essential.

Myth: Faster AI is used to monitor staff.

Fact: Responsible AI focuses on work support, not surveillance.

Myth: AI outputs are always right.

Fact: AI can make mistakes—staff review is required.

Myth: I need to be an expert to use AI.

Fact: Training is provided; tools should be simple and role-appropriate.

Myth: AI removes human connection.

Fact: When used well, AI frees time for people, not paperwork.

PEOPLE WITH DISABILITIES

WHY THIS MATTERS

AI should support your rights, choices, and voice—never replace them.

Myth: AI makes decisions about me.

Fact: You remain in control of decisions that affect your life.

Myth: AI means my voice matters less.

Fact: AI should support your voice, not replace it.

Myth: I don't need to understand AI.

Fact: You have the right to plain-language explanations.

Myth: AI means less human support.

Fact: AI should help staff spend more time with you.

Myth: Once AI is used, I can't opt out.

Fact: You should always have choices and the right to ask questions.

FAMILIES & CAREGIVERS

WHY THIS MATTERS

AI use affects trust, communication, and how services are delivered.

Myth: AI replaces staff decision-making.

Fact: People—not AI—remain responsible for care and support decisions.

Myth: AI means less personal care.

Fact: AI can reduce paperwork and increase time for relationships.

Myth: Data is automatically shared.

Fact: Data use is governed by policies, consent, and safeguards.

Myth: Families won't be informed.

Fact: Transparency and communication are essential to responsible use.

Myth: AI use can't be questioned.

Fact: Families can ask questions and raise concerns at any time.



NEW YORK
ALLIANCE FOR
INCLUSION & INNOVATION



Role Based AI DSP

- Line between supporting staff and changing the service relationship
- What does person centered AI mean to you
- What is the minimum capacity bar for someone to use AI
 - Training, consent, etc
- What decisions are made "by habit" that might quietly become "AI-assisted"
- What do you "measure" to ensure person centered choices stay true with AI use by staff



Where are DSP's using AI now

Documentation

- Supported through your EHR or not
 - Notes, monthly summaries, etc
- Reviewing current plans (life plan, safeguards, hab plans etc)

Planning support

- Meal ideas, recreational, crafts etc

What else?



Cultivate AI Culture



Space for questions, testing, uncertainty



Reinforce trust and not increase surveillance

AI supports your voice, your support for the person



Openly discuss the why

Important to you to save time and important for the people supported to have more of your focused time to live their best life.



Suggested Learning Sequence



NYAI Toolkit: Modules for DSP's

Module 1: understanding and communicating AI

- Glossary, understand what AI is, reduce fears and myths
- Gain confidence to engage without needing technical expertise

Module 2.4: Staff training and role based competencies

- Emphasis on human review, role boundary and practical application
- What DSPs can use AI for and where there's space for pause in using AI

Module 3: Risk, Ethics & Oversight

- Privacy, trust and person centered decision making
- Why not to enter PHI/PII into AI





NEXT STEPS

- What steps can/will your agency take to move forward with your AI journey
- Continued Role based conversations. Who at your agency is a department champion? What role next?





THANK YOU

Meetings will be on the 2nd Monday of
every month at 1pm

