

Establishing Your Own AI Policy and Protocols

2026 National Conference

ANDREW CLARK: The question we're answering here today is why do you need an AI policy? What we're really answering is beyond the obvious of any major processes or policy in our organizations. I'm sure we have them for all types of things, but AI, we have found, has just been a little bit more complex than that. It's more complex in terms of usage in governance. It's more complex in terms of people's perspective and how it interacts with their values. I'm sure there are a lot of very strong opinions about AI in this room alone, and it's very complex in terms of its ability to innovate. For a lot of our systems and software, you're going to have the software company itself telling you how to use it, what it can be used for, all of its perimeters and barriers. AI is a little bit more free-flowing than that.

So between it being a major new technological advancement that many people, even people who are a little more experienced with it, everyone is still struggling to understand the far limits of it, and it doesn't help that almost daily, if not weekly, there are major changes and updates to it. So a policy that really helps ground everyone in the allowable uses, how it interacts with your mission and your values, how it reacts to your culture, and how it might change your culture, and how you might be able to use it to innovate is really going to be crucial to have a policy to help you navigate that, because there are a lot of variables there.

We do just want to mention that there are many main types of AI, and for those who are a little less experienced, we just want to ground us all in about two of the main types. What we are really going to be talking about and what most people are familiar with is generative AI. Generative AI is when it takes one thing and creates something new that wasn't there before. So you put text in, you put a picture in. It brings you back everything. Those are the basics of ChatGPT, Claude, Google Gemini—all of the major ones that people use.

[00:02:00]

But what we are seeing increasingly in that, and some of you may have already started to work with and should be present in your policy, is around agentic or autonomous AI. And that's AI that can do things on its own, based off of a certain subset of data or instructions that you give it. So that's where it gets a little bit more into the idea of things happening without you, without human interaction, whereas most of generative AI is prompted, is directed by people. But for the most part, I would say probably most people here are more grounded in the generative AI than autonomous, but we did want to mention it in case any of you had started to use some like Claude Cowork or ChatGPT's Codex. That's kind of the next advancement. As you're putting your policy together, it is worth thinking about how you take on an AI that can do something on its own, without human prompting.

The way our session is structured today, we're first going to go through data protection and tool governance. That's really usually the very first question. Okay, we want to use it, but how safe is it? Are we able to put proprietary or sensitive information in there? We're going to address the importance of human oversight and create an integrity. Especially for our industry, it's very important for us to have a perspective on that. Again, we will not be very prescriptive on that here, but I think there are probably a lot of shared opinions about where AI can and cannot go into, when it comes to human creativity.

And finally we're going to talk about protocols and actually launching into policy. When you go back to your organizations this week, how do you go from a blank page to a full policy, and what are the steps that you can take? As I said, we'll be doing a lot of different table discussions. We left pen and paper there, and we would love to collect as much of that as possible. We had quite a lot of great orchestras reach out to us with their policy and kind of where they're at, but we're trying to get a better sense of just exactly where our members are and how we can be of better assistance.

[00:04:02]

So if you feel comfortable leaving anything behind, feel free to leave what you write down on the table, and we'll collect it, and hopefully it will help us with better resources in the future. So with that, we're going to go right into the first section, Data Protection and Tool Governance. I want to invite Nick up.

NICK COOPER: Good morning, all. Thank you for joining us today. I'm Nick Cooper. I'm the Executive Director of the Anchorage Symphony. Yes, Alaska exists, although not on every map, but we are here. I just want to start off—I was originally asked to do this, and I was like, "Sure," and then I met with both of these folks. The first thing I told them is, "I am not an expert. I use it a lot. I have used it a lot, but I am in no way an expert." So before I got to the Anchorage Symphony, a little bit of my background that will come into play a little later, I originally found myself in Alaska, working for the Division of Public Health. I was the Deputy Administrative Operations Manager, where I managed the COVID health response. So contracting, procurement, all of those kinds of things for the State of Alaska. After that, I was hired over into The Arc of Anchorage as the CFO, where I then found my way back into—I was a music teacher—back into my musical world.

So that's kind of the lens I'm bringing in through this, my multi-function areas that I've gone through. The first thing we're going to really talk about is what are you doing now? What's happening now in your organizations? I'm not going to read this slide word-for-word. You're going to have to do that.

[00:06:00]

But the concept is like, do you know what is happening in your organization? When I worked for the State of Alaska, there were hundreds of us working on the COVID response, and I didn't know who was using it when it first came out. I certainly was using it when it first came out, and nobody asked me anything about it. I just went on my merry way. So chances are, that may be the case in your organizations, where someone, somewhere is using AI that you may not know about. So giving that kind of baseline is incredibly important. Who is using it? Why are they using it? And then kind of going from there.

As you develop that and go through this discovery process, we want to talk about what risks you want to involve, and then this tension between this creating a policy while allowing for innovation. And so there is inherently this issue that lies where we don't want to lock things down immediately without letting it go along as it goes along. When I was first hired for the Anchorage Symphony, they asked me a question in my interview, and they said, "What's your vision for the symphony?" I don't know. I love the symphony, but I don't know the people. I don't really know the culture. I don't know the musicians. It's something that I have to figure out as I go along. And a lot of people, that concept of figuring out—because you do it as we did in Alaska for the COVID response. People are very uncomfortable with that concept. I think we need to be okay with that as we figure out, as we go along, how we make this happen. But put in those baselines.

[00:08:00]

A good example, I use a particular, I won't say which company, but it's a generative AI model. I started using it. I'd gotten really into it, and then suddenly I'm like, "I'm trying another one." I tried this other one. "Ooh, I kind of like this other one." And now I'm going back and forth and comparing them as I go. So again, allowing it to unfold as it goes along, I think you'll be okay with it. So along with the slide about what's happening in your organization, what's happening outside of your organization? Is your CRM using it? Is your media company using it? I can say we, our media company that we just got uses it. How do we deal with that? That's something to consider. Does that involve including it in a policy? A question to ask, yes?

I've got notes on top of notes here, so I want to make sure I get everything here. And then deciding what the comfort level that not only you have as leaders, but what your comfort level with your staff is. A staff member of mine outright refused, does not want to use it. And that's fine. I would encourage them to show use cases, if you are so interested in pursuing this. Use cases, because it doesn't have to be used maybe the way they think it would be used. They thought—and I'll use our case—they thought it had to be used to create images, or something like that. But no, we can use it to—we can plan out something. We can spitball

ideas with it. That's what I really use it for. So there are a lot of different use cases. So again, allowing that exploration to happen as we go along is really critical.

[00:10:01]

So we have a three-tiered approach here. I see lots of people taking pictures, and that's wonderful. This is a guideline. I'm not saying do this. I'm not saying don't do this. This is something you have to meditate on and come to your own conclusions that include your stakeholders. One thing I will very much suggest, however, again going back to my time working for the state is viewing this through the lens of HIPAA. We have a lot of HIPAA training, right? So that is the Health Insurance Portability and Accountability Act. This is the reason why doctors still use FAX machines, or why you have to log into a specific portal to email your doctor. It's a very conservative way of communicating. It also sets a minimum necessary standard as you provide data. A billing analyst or someone in the Finance Office doesn't necessarily need to know your entire medical history to be able to bill someone. They just need a little code to be able to bill someone. So there's disconnect, intentional disconnect between the doctor's side and the finance side.

And this is kind of the same thing. I think we should look through this as what kinds of information are we feeding into whatever type of AI we choose to use? There are some lower risks, some medium risk, and higher risks. What I think of it as is any kind of PII—and that's the HIPAA language—publicly identifiable information should probably be left off. So we don't include donor's names and amounts. It's probably not something we really want to put in there.

[00:12:04]

We want to assume that the stuff that we put in here is going to probably be publicly available at some point. I like to say don't put it in there if you don't want your mom to read it. Don't put it in there if you don't want your lawyer to read it. Don't put it in there if you don't want the newspaper to read it. Don't think of this as a super secure system. There's even a lot of talk right now amongst the legal community as to is this discoverable during the course of a lawsuit? I am not a lawyer. Talk to your own lawyers. That's all I'll say about that. But again, the minimum necessary standard, so contemplate that as when you're putting these items in there, as to how—what information it really does need to be able to function and do the processes that you'd like it to do.

RACHEL ROSSOS GALLANT: My name is Rachel Rossos Gallant. I'm the VP of Marketing and Membership for the League, and I'm going to talk a little bit about human oversight and creative integrity. This is something that's very important to me in my role and also for our field. I am also a musician, so I think about this in terms of my own creations, as well. We

started out by talking about this discovery process and protecting your data, but then also as you're developing a policy for your organization, it's important to have meaningful conversations around your organization's values and how they align with using AI tools. So here we have some examples of some of the tensions that may exist and how one might articulate this in a policy.

[00:14:09]

So artistic decisions, right? You don't want to outsource everything to a tool, but you can use it as a sounding board. Nick mentioned the idea of just spitballing things, so AI can inform, but it shouldn't be making choices about programming, repertoire, or any kind of creative direction. When it comes to creating materials for cultivation and relationship building, it can be a very helpful tool in writing, in getting you that first draft if you are stuck, or if you have your first draft, giving you that other perspective. But it shouldn't be doing all of the work for you. We're the ones who have that human touch. We're the ones who know what's going to be special to our audience members. We're the ones who can listen, and we're the ones with the nuance.

AI has a flattening effect on everything it does, the generative AI, because it is an aggregator. That's its job. It's looking at everything and learning from everything, and then it has a flattening effect. So we need to be very, very careful when we're using it to help us with creative projects. And then anything public-facing needs to go through rigorous review, as well, because it does hallucinate. It can introduce bias. So there are many different layers to that.

It's easy, also, when you think of AI versus people, and Andrew mentioned agentic AI being able to do more with fewer instructions and being a little bit more autonomous. It's easy to get caught up in the hype of AI replacing people.

[00:16:01]

There are some very real concerns about that, but that's not within the scope of what we, as an industry, where we are right now. Where we are is we need to figure out what these tools can do for us, how our people are already using them, and how we can put in the safeguards that we need to be able to use them safely and use them to innovate and find new possibilities. I'll also spotlight here some of the known AI limitations. Hallucination I mentioned already. Bias I mentioned already. Sycophancy is a big one. So AI is also—the LLMs, they are built to be very agreeable and to reinforce opinions that you may already have. It's similar to social media algorithms that are going to just feed you more and more things that will get you to click that "like" button or comment on it. The more agreeable it is, the more likely you are to keep using it. So you just need to be very, very mindful if you are

digging into these tools, that that is a real risk. You can tell it to be disagreeable, but it's not going to do it on its own.

So you can say, "What am I missing here?" Or, "What are some of the other perspectives?" But you have to actually prompt it to do that. So in all cases, AI output is our starting point. It's not the final product. And Nick is going to give us a few real world examples of how he has been using it.

COOPER: So just to put this into context, I use AI really as a sounding board and to ask questions and probe and ask more questions.

[00:18:04]

That's really where I found its greatest use, and to just brainstorm. That's how I operate as a manager. I'm always looking at my staff, "Hey, what do you think?" Or board members, "Hey, what do you think?" That's just how I operate. It makes me think better. This is great for that. One thing I did do successfully with it after a while was really use some real-world data and decide how I am going to rezone and re-price our hall. And I took statistics in graduate school, and I could probably have programmed this into R over a month and come up with these answers. Here I was able to do it in—I mean, it probably still took a month, but partial days of working on it. But there was a lot of back-and-forth. There was a lot of restarting these conversations. So it's very much a don't accept what you get the first time around.

I'll give you another example, and I forgot my prop. I was going to bring it, but my bio that I wrote for this. If you go on the old thing, you'll see it. The conversation it took me to get to the bio where I liked it was 27 pages long when I went to print it out. So I am very particular. So, yes, I think it does save time in a certain sense. Does that mean that no time will be spent on something? No. As Rachel said with the last slide, we have to have this overview of the human overview of it. So again, especially the more nuanced the thing is, whatever you're working on—again, when I was using all of my sales velocity data to get at the price elasticity of demand and all of that good stuff, it just was a lot.

[00:20:05]

This was also a few months ago, I will say. This was back in January. Who knows how fast it would go now? It is a wildly different world now than it was back in January. So that's another thing to consider, as well. Also, sorry. We've had to cut things a little short on our conversation because I talk too much. I should have warned them that giving me a microphone is not a great idea. So with that, I did have a place where I did something, I got to a good spot, and I said, "Yes, this is great." And it was with a refreshed employee handbook. And I'm like, "Great, we have a 4-person organization, plus some part-time folks. This is a

good spot, and I like it. I'm going to send it to my buddy, who's a real HR guy, the Director of HR at [UNINTEL] Hospital. And he goes, "No. No, sir."

It was 20 pages, maybe. He returned it back, and it was 35. So having that kind of real expert review stuff that real experts need to be involved in, that really needs to be contemplated. I'm pretty dangerous, because I know a little bit about everything, but I recognize that. So I am able to reach out in the areas that I do need to reach out. So just something to consider, right? You need that person who is going to catch what the AI gets wrong.

ROSSOS GALLANT: Thank you, and I will add actually another example to that. One of the things that can be really tricky about using AI as a tool is that you can end up missing out on knowledge transfer because if somebody puts together a document, and then someone else runs it through AI but doesn't actually read through it, then things can easily get lost in translation.

[00:22:03]

I know for me, I learn better if I'm reading. So I always make a point of reading things before using my AI tools to help out. One example of that is the brand new Governance Guide that just came out. When I was editing that, I read that from top to bottom before I reached out to AI to help me find things that I might have missed. And that was really helpful because then, when it gave me a list, I knew where it was giving me real information and where it was something that I could just let go. There was one point—I really only brought it in late in the process, where I was cross-checking. Here was our long list of 50 very small edits that we had given to the designer. Please tell me: Which of these can you tell has already been applied to the PDF? And it gave me back a nice list that had green, yellow, and red. For green, yes, it was done; yellow, I can't check this because I only look at texts; and then red, this needs to be addressed. But even with that, there was one that was fixed, and it was text. And it kept saying, "This is wrong. This is wrong." And I said, "Well, tell me more about this." And eventually, with enough prompts, it said, "Oh, this is correct now." So nothing had changed in the PDF I was processing, so it's just something to keep in mind.

Just a last thought before we go into table discussions is coming back to creating the policy. We talked about discovery and finding out who's using it, how they're using it, what some use cases might be, that people are on very different ends of the spectrum as far as their comfort level and their enthusiasm for it.

[00:24:04]

And so a really important part of getting buy-in behind a policy is keeping people engaged in the entire process. So you need somebody who can really own that process from top to

bottom. You need some collaboration. We suggest forming a committee. Maybe that's too formal for your organization, and that's fine, but the idea just being that you're identifying your stakeholders and making sure the right people are in the room, and then keeping that dialogue going all throughout the process, and then finally establishing that oversight in governance. Andrew will talk more about that in the next section.

CLARK: Okay, so now we're going to get into really what we said this session was about. The first two pieces are the crucial pieces of background that you actually have to consider before you even really get into making your plan. You have to understand your culture. You have to understand data security. You have to understand what goes into AI before you can really start writing it. But from there, what we're really talking about is how you take that and now turn it into a policy. And because AI is so sensitive, because it has such a wide range of things, there are a few ways that we think that you can translate those values into protocols.

So defining both opportunities and risks. Anyone who has talked about AI in all sunny or all negative terms is not being honest about it. You need to think about all the positives and all the negatives as you go into it, so that you can create a balanced protocol and policy that people are able to accept. If you come into the room and say, "This is the best thing ever. We're going to make all of our lives easier. This is going to be simple, simple, simple," then they're not going to accept that because they know that there's more to it. Making sure that you align it with your values. So you've talked to people. You've talked to your stakeholders. You've discussed it amongst your staff, what's important. If those conversations don't turn into something in the policy, then there's not going to be buy-in. They're not going to be interested. They're not going to feel heard.

[00:26:11]

So as you have these conversations, even at the beginning, think about in the moment what's coming out of those conversations that can help create something new. This is not a simple policy of, "Here's how you put in a record in Salesforce." There are a lot of larger questions when it comes to AI that people are looking to see in their policies so they can feel comfortable. And then these bottom two are kind of where you can really go beyond, "Here are the nuts and bolts; here's a policy," but also a way of being with AI that you can go through. Really spend your time researching which tool is best for you.

I'm sure you all have talked about what each other are using today. Everyone has different needs. Everyone has different opinions. Someone just likes different interfaces more. I would really like the time. Don't just go with whatever is easiest. Don't even go with whatever is cheapest because those prices are going to change in the years to come, and you're on a monthly solution that's best for you. Nick had mentioned this earlier. Something that's really, really helpful with AI that I've not really experienced as much with other

software is use cases. In a lot of ways AI is only as helpful as your imagination. It can do so many things. The more people in the room who talk about what they're doing and how they're doing it, the more it opens up everyone else's eyes. I would even recommend finding some, whether it be formal or informal, some way to create a use case library that really kind of circles through, "Here's where I've had success. Here's where it didn't work the way that we talked about." Some people are going to go into AI and say, "I know how to write an email. I'm not needing you to correct my email. I know how to write a press release. I can do that." There's a lot of the basic AI stuff that people who are not getting into it are kind of confining what AI can do.

[00:28:08]

And the more you say, "This is this basically grunt work that I have to do as part of my job, that takes 10 hours a month." I can get this down to 2 hours, and now I can maybe focus more on the creative work instead of all the administrative work that goes with so many of our jobs. That's where people really start to get a little more excited. I think we all know no one's running AI to take away the work that they love to do, but I would be impressed if anyone here doesn't have a bunch of side work that they have to do to keep their job going, and that's really where a great place for use cases comes in.

And finally, the other big piece here is around training and periodic review. The good thing about AI when it comes to training is if you're feeling overwhelmed, if you're feeling like it's too early, AI is really the only software that can help you learn how to use it. You can ask it, "Can you do this for me?" And it might say, "No, I can't, and here's why." Or it might say, "Absolutely, here you go." And then you can play with it from there. Certainly most software is not that way, but even with that, if you are not providing training to your staff, even if it's informal, even if it's free training or webinars. It could be YouTube videos, just something that walks you through how to use these tools, you're going to have a policy that goes nowhere. If you're feeling overwhelmed, if you've taken the time to put a policy together, and you just hand it out to everyone and say, "Now you can use AI," it's just not going to go anywhere, because they need to understand how to use it.

So really good training. We even have a few business partners here at the conference that help with this. Training to make sure that the policy goes somewhere. But as I mentioned at the top of the presentation, one of the most complex things about going through AI is actually how quickly it moves.

[00:30:01]

Nick had mentioned that what was AI in January is not what AI is today. When we started putting this policy together, it was a completely different landscape. And that's really

difficult to keep up to date with, and the only real way to do it is to try and stay informed and give yourself every three months, every six months, maybe more, to look at your policy and see: Is it addressing the needs? Do we feel comfortable with how everyone is using it? Does everyone feel informed about what they're allowed to do, versus what they're not allowed to do with AI? So this is going to have to be a living, breathing document, maybe even closer to guidelines than probably most of your policies are, but it will give you a world of opportunity if you can actually keep that dialogue open with all of your staff.

I want to touch a little bit more on the training itself, because one of the reasons it's important, beyond just that people feel comfortable, is risk to the organization. AI's uses are feeling limitless right now. There are a lot of different approaches. People are doing incredible, incredible, huge projects with it, but if you do not train them properly, do not make sure that they're on an enterprise system, if they're putting proprietary information in, if you're not making sure that the data is secure and that people are using the tool that you've said they're allowed to use, you could find that all your proprietary information is out there, that donor information is out there, that staff information is out there.

So it's very important to be clear. We want all of our staff to be experimenting, trying different things with AI, but they need to understand the guidelines, and they need to be trained to do so. Otherwise, you will find yourself in a position where you've allowed very sensitive information to go out.

Everyone is going to come to their own versions of their policies. We are actually going to distribute our own policy, and by the way, we'll also be distributing this.

[00:32:04]

Sorry, I should have said before, when everyone was taking pictures, the whole outside will be distributing these slides, as well. But we got quite a lot of policies from our members, as we start to prepare for our own for this session. The most common ones are the ones here, basically all the things that we've talked about throughout this session. What is AI, the definitions, the scope of how you're going to use it? What are the expectations of all your stakeholders? What's allowed, versus what's prohibited? It's such a wide range, you're never going to hit all of it, so again think guidelines more than specific uses. What is required of all staff before they start using it? That's also going to fall back on you to make sure that they're being trained. Which tools are they allowed to use? This could even be AI tools within software that you pay for and exploring all the different AI capabilities there.

And then finally, you're really going to want to think about governance and oversight. This is, in my opinion, one of the harder pieces for AI because everyone is in their own instance, doing their own thing. You can put certain limitations on it based off the tool you're using,

but it's not as easy to see how everyone is using it. So again, training, training, training. Continue the dialogue. Make sure everyone is aware of what they're able to do. And when in doubt, make sure that if you're able to, pay for a service that does not train the model based off your information. Enterprise levels are the best to make sure that at least they have a little bit of a more soft landing pad for all the experimentation they're doing.

So to kind of round us out in terms of everything that we've covered here—we've gone through all these different pieces and all these different things that we recommend that you do as you put your policy together. I'm just going to go through them all again very quickly to round us out. One is gathering input, talking to whomever it is who would be interested in talking about it—definitely your staff, likely your board, potentially even other stakeholders, if you find that they would have a strong opinion.

[00:34:08]

It is definitely on its way. Most people are using AI at this point, but there are very strong opinions, and it could affect other relationships you have.

The next one is diving into your first draft because there are so many fears and a lot of times excitement about AI. We recommend that you dive in with your first drafts so people have something to look at. I think just by everyone attending this session here, you can tell that a lot of people do not have a strong sense just yet of what direction they want to go in with AI. Looking at something first, I personally have found it to be a lot easier than testing every single idea with a large group of people.

From there, you're going to want to get all your group buy-in, get everyone onboard as much as possible. As Nick mentioned, we've noticed that most staff have a few people who are holdouts, who really don't want to come onboard. For the most part, you want to have large buy-in, hopefully large excitement, widespread training and policy review. And then set everyone off into doing it. Beyond the use cases, beyond training, most of this is going to be a largely personal experience. Everyone has to think for themselves, "Where do I want to off-load some of my work to AI? What can I test to do?" And the most you can do is keep the communications open, keep the group discussions going, because it's really going to be most helpful to have a group learning activity in order to get everyone to get on the same page about AI, because it's going to work completely differently for every single person, which is a little more challenging than some other software.

[00:35:54]