

**CONSEJO NACIONAL DE ACREDITACIÓN
ACTA DE LA SESIÓN ORDINARIA 1977-2026**

SESIÓN ORDINARIA DEL CONSEJO NACIONAL DE ACREDITACIÓN DEL SINAES
CELEBRADA EL PRIMERO DE SETIEMBRE DE 2026. EN LAS INSTALACIONES DEL
SINAES. SE INICIA LA SESIÓN A LAS NUEVE Y DIEZ MINUTOS DE LA MAÑANA.

MIEMBROS ASISTENTES

Dra. Susan Francis Salazar, presidenta
Dra. Lady Meléndez Rodríguez
Dr. Ronald Álvarez González
Ing. Walter Bolaños Quesada

Dr. Álvaro Mora Espinoza, vicepresidente
Dra. María Eugenia Venegas Renauld
M.Sc. Gerardo Mirabelli Biamonte
MAE. Marianela Núñez Piedra

INVITADAS HABITUALES ASISTENTES

PhD. Melania Brenes Monge, directora ejecutiva.

M.Sc. Heilyn Vásquez Hernández, asesora legal.

Mag. Marchessi Bogantes Fallas, coordinadora de la Secretaría del Consejo Nacional de Acreditación.

***Los asistentes se encuentran conectados virtualmente.**

Artículo 1. Revisión y aprobación de la propuesta de agenda 1977.

Dra. Susan Francis Salazar:

Buenos días a todos y todas. Estamos hoy en sesión híbrida, 1 de setiembre del 2026, sesión 1977. Ustedes cuentan con la agenda, pudieron dar cuenta de que hay dos puntos que son los de fondo, hay un punto de informes que se coloca nada más por una formalidad en términos de lo que está reglamentado para este Consejo. Si no hay observaciones con respecto a la agenda, entonces sírvanse indicar de una vez su aprobación. Muchísimas gracias.

Para el día de hoy no contamos con acta, dado que el día de ayer fue feriado y no hubo tiempo para prepararlo.

Se aprueba agenda 1976 de manera unánime.

1. Revisión y aprobación de la propuesta de agenda 1977.
2. Informes.
3. Análisis focalizado del Informe de Comisión de Pertinencia 2025.
4. Propuesta estructura ponencia-foro INQAAHE 2027.

Temas tratados: 1. Revisión y aprobación de la propuesta de agenda 1977. 2. Informes. 3. Votación de declaración confidencial. 4. Análisis focalizado del Informe de Comisión de Pertinencia 2025. 5. Propuesta estructura ponencia-foro INQAAHE 2027.

Artículo 2. Informes.

De la Presidencia:

Dra. Susan Francis Salazar:

Procedo con los informes de presidencia. Solamente tengo un punto que es que estamos formalmente ya invitados para la reunión con SIACES, una reunión que es aparte, ya están las dos, es una del Comité Ejecutivo y la otra donde los miembros del SIACES quieren conocernos a don Álvaro y a mi persona. Entonces, ya está formalmente establecida. De ahí, imagino yo que habrá algunos puntos que ellos quieren trabajar sobre la solicitud que se hizo para que Costa Rica fuese anfitrión, entonces probablemente sería como el punto a reseñar en ese espacio. Ese sería el único punto de informe que tengo yo en este momento.

De los Miembros:

Dra. Susan Francis Salazar:

¿Los miembros tienen algún informe? Adelante, doña María Eugenia.

Dra. María Eugenia Venegas Renauld:

Buenos días. Quisiera aprovechar el inicio de este mes que tradicionalmente ha sido dedicado a celebrar y festejar la independencia del país, para compartir con este Consejo la importancia de mirarnos en una circunstancia particularmente delicada que nos retrotrae de conquista, de libertad y desarrollo nacional como nunca antes. En particular, cuando

como civilización contamos con tanto desarrollo tecnológico, tanto conocimiento y tanta información, y resulta hasta desolador, más bien ir como en sentido contrario a lo que la humanidad ha hecho y ubicarnos en épocas muy oscuras. Entonces, creo que es un mes que nos ofrece la posibilidad en espacios personales, de hogar y también organizacionales para mirar lo que nosotros como país tenemos, eso es parte también de analizar la calidad de los servicios sociales, en este caso de la educación superior, para mirar realmente todo aquello en lo que nosotros podemos más bien consolidar. Y creo que el tema que nos convoca hoy es una buena metáfora como para repensar cuál es el papel que tiene una organización como en SINAES que busca y tiene por mandato verificar la calidad en un contexto tan convulso del mundo y no solamente nacional. Eso era lo que quería compartir con ustedes. Gracias.

Dra. Susan Francis Salazar:

Muchísimas gracias, doña María Eugenia, por la reflexión cívica, sobre todo subrayando el papel que como institución del sector público asociada a la educación como valor social nos corresponde. Doña Lady.

Dra. Lady Meléndez Rodríguez:

Buenos días a todas y todos. Nada más es recordar, Susan, que posiblemente en esa reunión de SIACES también estén haciendo el traslado del puesto que ocupé yo dentro de la junta directiva, que entonces se hereda a la próxima presidencia. Entonces, no sé si habrá alguna juramentación o algo por el estilo como un nuevo miembro de la junta de SIACES.

Dra. Susan Francis Salazar:

Ok. Muchas gracias.

De la Dirección:

Dra. Susan Francis Salazar:

Continuamos con algún informe de la Dirección Ejecutiva.

PhD. Melania Brenes Monge:

De mi parte, hoy no tengo informes.

De la Secretaría:

Dra. Susan Francis Salazar:

De la secretaría.

Mag. Marchessi Bogantes Fallas:

Buenos días, no tengo informes.

Dra. Susan Francis Salazar:

De acuerdo.

El M.Sc. Gerardo Mirabelli Biamonte, ingresa a las 9:15 a.m.

Artículo 3. Votación de declaración confidencial.

Continuando con el desarrollo de la agenda, este Consejo procede a valorar la necesidad de declarar confidencial el siguiente punto, en atención a la naturaleza de la información sometida a conocimiento del órgano colegiado.

CONSIDERANDO:

1. Que la información, documentación y deliberaciones vinculadas al presente asunto contienen insumos técnicos, criterios jurídicos, elementos de análisis institucional y actuaciones administrativas en curso, cuya divulgación podría afectar el adecuado desarrollo de los procedimientos administrativos, la toma de decisiones del órgano colegiado y la eventual determinación de responsabilidades.
2. Que el artículo 24 de la Constitución Política de la República de Costa Rica tutela el derecho a la intimidad, así como la protección de la información de carácter sensible o confidencial, permitiendo limitar el acceso a aquella información cuya divulgación pueda comprometer intereses públicos o derechos de terceros.

3. Que corresponde al Consejo Nacional de Acreditación resguardar la confidencialidad de la información que, por su naturaleza, requiera un tratamiento restringido, garantizando con ello la integridad del proceso administrativo y la protección de los derechos de las personas involucradas.

SE ACUERDA:

1. Declarar de carácter confidencial el punto de agenda relativo al “Análisis focalizado del Informe de Comisión de Pertinencia 2025”, conocido en la presente sesión, al amparo de lo dispuesto en el artículo 24 de la Constitución Política de la República de Costa Rica.
2. Disponer que en el acta respectiva no conste la transcripción literal de la información presentada ni de las deliberaciones efectuadas sobre este asunto, dejando únicamente constancia de la presente declaratoria de confidencialidad y de los acuerdos adoptados.
3. Establecer que la documentación relacionada con este asunto deberá resguardarse bajo las medidas administrativas correspondientes, limitando su acceso exclusivamente a las personas funcionarias y miembros del Consejo que, en razón de sus competencias y funciones, deban conocerla.
4. Indicar que la presente declaratoria de confidencialidad se mantendrá vigente mientras subsistan las circunstancias que justifican la restricción de acceso a la información, o hasta tanto el Consejo disponga lo contrario mediante acuerdo expreso.

Votación: 7 votos a favor y 1 voto en contra por parte de la Dra. Lady Meléndez Rodríguez.

Dra. Lady Meléndez Rodríguez:

Yo creo que el hecho de estar haciendo todo un esfuerzo importante por hacer una transformación hacia la mejora de un sistema al cual nosotros pertenecemos, porque pertenecemos a un sistema, debería ser de conocimiento que estamos haciendo estas acciones hacia mejoras y hacia una transformación. Me parece que lo que estamos discutiendo es precisamente los preliminares de esos grandes cambios y que posiblemente los otros componentes del sistema están previniendo. Me parece que todo esto perfectamente puede ser de orden público. Gracias.

Artículo 4. Análisis focalizado del Informe de Comisión de Pertinencia 2025.

Los miembros del Consejo realizan el análisis focalizado del Informe de Comisión de Pertinencia 2025.

RESULTANDO:

1. Que de conformidad con la Ley N.º 8256 y sus reformas, corresponde al Consejo Nacional de Acreditación dirigir y definir las políticas, lineamientos y normativa general que orientan el quehacer institucional del SINAES.
2. Que, en la sesión N.º 1828, celebrada el 7 de febrero de 2025, se conformó la comisión de trabajo encargada de la reflexión y elaboración de una propuesta de reforma del Reglamento Orgánico, iniciativa que posteriormente fue objeto de modificaciones en las sesiones N.º 1830, del 14 de febrero de 2025, y N.º 1884, del 5 de setiembre de 2025.
3. La Comisión de Pertinencia de la Organización (CPO), conformada por acuerdo del Consejo Nacional de Acreditación (CNA), presentó en la sesión N.º 1937 del 10 de abril de 2026 el Informe de la Comisión de Pertinencia de la Organización 2025, el cual contiene el análisis efectuado sobre la pertinencia organizacional del SINAES, así como un conjunto de conclusiones y recomendaciones dirigidas al fortalecimiento institucional y a la modernización de la organización.
4. Que en la sesión 1961 del 7 de julio de 2026, la Dirección Ejecutiva y la Unidad de Asuntos Jurídicos elaboraron un documento en el que se identifican las principales acciones proyectadas por la Administración en atención a las recomendaciones formuladas por la Comisión de Pertinencia de la Organización, con el propósito de facilitar el seguimiento institucional y aportar elementos para el análisis del Consejo Nacional de Acreditación.

CONSIDERANDO:

1. Que el Informe de la Comisión de Pertinencia de la Organización 2025 constituye un insumo para la reflexión del Consejo Nacional de Acreditación sobre las condiciones actuales de la organización, sus necesidades de fortalecimiento y las oportunidades de mejora institucional.
2. Que los miembros del Consejo Nacional de Acreditación analizaron de manera focalizada los modelos SINAES percibido y sugeridos definidos en el informe presentado por la Comisión de Pertinencia de la Organización, lo cual fue objeto de discusión durante la sesión en que fue presentado.
3. Que, como resultado del análisis efectuado, el Consejo Nacional de Acreditación considera pertinente utilizar los insumos derivados del Informe de la Comisión de Pertinencia de la Organización como base para orientar una fase preliminar de planificación estratégica, en articulación con los procesos institucionales que se encuentran en desarrollo.
4. Que dicha fase preliminar requiere considerar los elementos identificados por el Consejo Nacional de Acreditación, incluyendo la revisión del ser, la misión y visión institucionales, la definición de un escenario futuro, las prioridades estratégicas, las necesidades operativas y los avances institucionales existentes.
5. Que el proceso de planificación estratégica debe mantener al Consejo Nacional de Acreditación como instancia responsable de la conducción y definición de las orientaciones estratégicas institucionales, con el acompañamiento técnico que corresponda por parte de personas asesoras, facilitadoras o especialistas en planificación.
6. Que resulta conveniente que la metodología que se adopte permita establecer resultados, metas e indicadores diferenciados según distintos horizontes temporales, de manera que facilite el seguimiento y evaluación de los avances institucionales.

SE ACUERDA:

1. Dar por recibido el Informe de la Comisión de Pertinencia de la Organización del Consejo Nacional de Acreditación del SINAES, correspondiente al año 2025, como documento base para las decisiones estratégicas posteriores y para el desarrollo de una fase preliminar de planificación estratégica institucional.
2. Acordar, como paso inicial del proceso de planificación estratégica, el desarrollo de una fase preliminar orientada a armonizar e integrar los insumos derivados del Informe de la Comisión de Pertinencia de la Organización y los avances institucionales existentes, como base para la posterior formulación del plan estratégico.
3. Solicitar a la Dirección Ejecutiva la aplicación de una propuesta metodológica para el desarrollo de esta fase preliminar, que permita avanzar progresivamente en la formulación del plan estratégico e incorporar los avances institucionales existentes, las acciones emprendidas en atención a las recomendaciones de la Comisión de Pertinencia de la Organización y los avances en la actualización del Reglamento Orgánico, en lo que corresponda y resulte viable a la fecha. La propuesta deberá ser presentada al Consejo Nacional de Acreditación en la última sesión del mes de noviembre de 2026.

Votación unánime.

Artículo 5. Propuesta estructura ponencia-foro INQAAHE 2027.**Dra. Susan Francis Salazar:**

Damos paso al último punto, que es una propuesta de estructura para el foro INQAAHE 2027, que en este caso yo no solo le agradezco, sino que le daría la palabra a don Ronald, quien nos la compartió para que él se refiera y nosotros podamos estudiar y analizar cuál es su propuesta. Adelante, don Ronald.

Dr. Ronald Álvarez González:

Es algo que está conectado con lo que estábamos hablando sobre la propuesta de transformación digital. Voy a compartir aquí la pantalla para compartir la presentación.

Traté de hacer una propuesta basada en una serie de hallazgos que yo he encontrado a lo largo de diferentes propuestas que se han hecho al Consejo y al SINAES, en términos de cómo terminamos de definir la aplicación de una tecnología que es emergente, que está en sus fases de modelación, pero que no se termina de construir todavía pero que tiene un impacto tremendamente actual y disruptivo en lo que significa el uso de esta herramienta que es la inteligencia artificial. Hice una serie de investigaciones sobre cómo aterrizamos esto en el SINAES. Después de tener una charla con Melania y los muchachos de TI, llegué a la conclusión de que ellos mismos les cuesta pensar en cómo aterrizar esto. Entonces, hay cinco elementos que ustedes vieron ahí en la presentación, elementos que yo considero, no solo yo considero, sino que la literatura técnica los avala y que me parece que son fundamentales. Uno puede empezar hablar de por cualquier lado, pero yo prefiero contarles que cuando voy a caminar ahí por el campus de la UCR me encontré un banner gigantesco que me encanta, el banner está en el edificio de Estudios Generales y el banner es muy sencillo, pero es gigantesco y dice: Humanismo que construye, pensamiento que transforma y diálogo que conecta. A mí, me gusta porque justamente eso es el mayor reto que tenemos hoy en día para adoptar la inteligencia artificial y es el centro de discusión de todos los que pensamos que adoptar la inteligencia artificial no es algo que sea tan sencillo si no nos cuestionamos realmente para qué nos sirve eso. Por eso fue que yo les mandé ese artículo de Carmody Gray y que después me di a la tarea de ubicarlo, porque a mí me parece que las reflexiones que esta filósofa hace son realmente lo que nos debe de cuestionar. Ella dice que la invitaron los proveedores de inteligencia artificial a participar en una conversación sobre para qué le puede servir esto a la humanidad y resulta que cuando ella llega a conversar con los proveedores de inteligencia, con Anthropic específicamente no eran todos, era específicamente Anthropic, se da cuenta de que en realidad no era ese el foco de la conversación, lo que ellos querían era otra cosa, lo que ellos querían era ver cómo convertían a la inteligencia artificial en un ente pensante y un ente con conciencia.

Y ella dice, "oh, aquí hay una gran diferencia, esto no es lo que yo esperaba de la conversación". Esa no es la pregunta correcta, la pregunta correcta es que las computadoras no piensan, que piensa el ser humano con las computadoras, pero no al revés. Y entonces ahí es donde hace todo un análisis de la famosa encíclica de León XIV sobre esta famosa y muy comentada en todo lado encíclica sobre el uso de la inteligencia artificial y las implicaciones que tiene para la humanidad. Y ella dice, el verdadero aporte de la encíclica es que si hace las preguntas correctas y saca el foco del análisis que los proveedores de inteligencia artificial quieren tener. Por eso están contratando filósofos, porque están intentando convertir a la inteligencia artificial en un ente pensante per se.

Es interesante eso, porque me conecta con dos cosas más. Una, es que en realidad según el titular de la edición del domingo de hace como tres semanas de la revista Economist decía, tenemos que ver a los proveedores de inteligencia artificial como cuidadores de animales peligrosos, eso decía el título. Y cuando uno empieza a leer el contenido, se da cuenta de que no son solo peligrosos los animales que cuidan, sino que son peligrosos los cuidadores.

Y dice Carmody Gray que por más que conversó con ellos, que le parecieron simpáticos, no puede uno abstraerse del hecho de que lo que ellos quieren es vender un producto. Entonces, resulta que sí nos quieren vender un producto y yo les cuento que yo uso inteligencia artificial y al final de un día de trabajo ese producto con el que yo estoy conversando casi como si fuera una persona, al final me dice, "bueno, perfecto, hemos logrado a,b,c,d, ha sido muy productivo. Te espero mañana. Un abrazo". O sea, se transforma como en una persona, han tenido la habilidad de darnos un producto que en el fondo nos está engañando, porque como dice Carmody Gray, la relación va en desventaja.

¿Por qué en desventaja? Porque, en el fondo, la plataforma sabe más de uno que lo que uno sabe de la plataforma, la plataforma sabe quién es uno, sabe cuáles son sus gustos, sabe cuáles son sus intereses, pero uno no sabe nada de la plataforma, es una caja cerrada.

Entonces, desde esa perspectiva el tema de ese cuarto aspecto que aparece ahí de Human in the loop como humano céntrico se vuelve fundamental, pero se vuelve fundamental, no solo para decirlo como una declaración general, no. Cuando fuimos a visitar a una universidad que nos invitó, ellos hablaban de modelos humano-céntricos, pero el problema es que todo mundo dice eso, sí, pero ¿cómo lo hacemos?

Cómo lo implementamos para que realmente ese modelo HITL, Human in the loop, realmente opere como tal, garantice que realmente funciona.

Y ahí les tengo que contar que lo que nos envió la Universidad de Costa Rica, esos cinco archivos que nos envió, en dónde tengo que contarles que se supone que nosotros estábamos invitados, no participamos, me consultaron una vez, yo les envié un documento, pero ellos lo trabajaron. El resultado de esos cinco documentos es fantástico, me encantó, es excelente y eso es lo importante, o sea, lo que ellos hicieron a mí me parece que es de muy alta utilidad y me parece que es una referencia que vamos a tener que tomar en cuenta en nuestras políticas de adopción de inteligencia artificial. Yo lo estudié con mucho cuidado, ayer lo estuve viendo con mucho cuidado, además me animé a hacer un anexo adicional, o sea, la versión que ustedes tienen del informe que yo presenté, tiene dos anexos más ahora.

Una es la evaluación de la propuesta desde la perspectiva, después de haberla leído con cuidado de la propuesta que hace en la Universidad de Costa Rica sobre cómo debemos de adoptar las políticas de inteligencia artificial y ahora les voy a contar, el cierre va a ser ese, cómo queda el análisis de la propuesta nuestra, porque quiero decirles que mi sueño es que sea nuestra y no mi propuesta, cómo queda ese análisis a la luz de este planteamiento que hace la Universidad de Costa Rica.

El segundo anexo, porque ustedes solo tienen el primero, y no se los envié antes porque es que eso lo trabajé ayer. Era adelantándome a una posible pregunta de ustedes y si no la han pensado, se las adelanto. La pregunta que yo me imaginé es, ¿en la experiencia real, hay alguien desde el punto de vista de las agencias de acreditación, de las instituciones académicas de universidades, de instituciones públicas, incluso del sector privado que hayan hecho algo parecido? Yo me imaginé que esa podía ser una buena pregunta, entonces me adelanté y sí hice una investigación rápida de fuentes y ahí quedó como un anexo.

Dicho todo eso, pensando que ese es, aunque tal vez no está en el documento presentado en ese orden, porque el primer está en ese orden que está en esa filmina.

Hay cinco elementos que son esos que yo llegué a identificar por la literatura, que son aquellos que tenemos que tomar en cuenta para hacer una presentación, y aquí sí confieso que es un asunto de aterrizaje de la adopción de la inteligencia artificial de manera adecuada para el SINAES y cuáles son esos elementos que tenemos que tomar en cuenta. Y son esos que están ahí, ya iremos viendo muy rápidamente, ustedes lo leyeron, así que no tengo que ahondar mucho.

Para ir en orden, lo primero que aparece es la gestión autónoma mediante herramientas agénticas ¿Qué es ese cuento del uso de agentes? Bueno, que la inteligencia artificial hoy en día nos presenta una herramienta práctica que tiene una serie de implicaciones que tenemos que de alguna manera tomar en cuenta, porque la inteligencia artificial ya avanzó hace mucho de ser un simple chat de consulta a una aplicación específica para lograr objetivos concretos. Por eso, es que ahí aparece que en el contexto que el SINAES es un garante de la fe pública, su misión fundamental es asegurar la calidad de la educación superior y las leyes del SINAES, que tiene que haber un puente que es el plan estratégico

sobre lo que hemos hablado mucho y que ese puente lo tenemos que construir, todavía no lo hemos construido, y que tenemos que tener una visión específicamente sobre la adopción de la inteligencia artificial, una visión hacia la transición de la automatización inteligente hacia la autonomía agéntica. Dice ahí abajo, la inteligencia artificial ha dejado de ser un loro estocástico, como se le dice ahora, que repite, que te dice, “un abrazo, mañana continuamos”. ¿Por qué estocástico? Porque es aleatoriamente la mejor respuesta posible y eso es lo que hace un chatbot. Y se convierte en un esquema en donde una organización como el SINAES puede lograr estructurar tareas específicas.

Hay un espectro agéntico que yo traté de ilustrar con ejemplos posibles dentro del SINAES, no me voy a poner a ver los detalles, pero hay cuatro posibles agentes que nosotros podemos incorporar en un modelo nuestro, uno reactivo, uno deliberativo, uno anticipativo y uno prescriptivo. Yo puse ahí algunos ejemplos de cada uno de ellos.

Hay una disyuntiva bien interesante en esto que son dos posibles estrategias que uno puede adoptar y que casi que se cae de la mata, como dicen, que nuestro enfoque debe de ir hacia la segunda, hacia el enfoque de que se llama el Capacity-up y no el enfoque Cost-out, porque el enfoque Cost-out lo que significa es sustituir personas con los agentes, que obviamente es posible, pero que no es eso lo que queremos si partimos del enfoque Human in the loop, obviamente. Y lo que queremos es llegar hacia ese enfoque de Capacity-up o la organización que ahora se le llama Centauro, ¿por qué? Porque es el ser humano potenciado con las herramientas agénticas. Ahí el objetivo, expansión de las capacidades mediante la colaboración simbiótica humano-agente. Ese es un término que se le da en la literatura especializada de tecnología a ese enfoque, no caer en la trampa de Turing, ese es el asunto, la trampa de Turing es justamente lo contrario.

Ahí caemos en un segundo punto que es fundamental, que es el uso de la IA como metalenguaje. En donde ya no es solo le pregunto al loro estocástico, y el loro estocástico me responde, no, caemos en algo todavía superior a lo que se dio en llamar ingeniería de prompts, que lo que nos decían, era, “mire, usted diseñe el prompt porque lo que usted le diga a la IA de una manera muy acotada, así va a ser la respuesta. Si usted le hace una mala pregunta, le va a dar una mala respuesta”. Resulta que ya eso evolucionó, ahora ya no es solo hacer un buen prompt, no es que usted puede hacer metaprompting, ¿qué significa eso? Que usted puede potenciar el diseño de sus prompts en un esquema técnico. Yo les cuento que en experiencias que yo he tenido, yo utilizo una inteligencia artificial para ser auditor de otra inteligencia artificial y a veces caigo en círculos que no me puede resolver lo que quiero y la otra me hace una auditoría de lo que el otro está programando, de y estoy hablando de generación de código. Y entonces me dice, “ok, aquí está el problema, no te lo está pudiendo resolver porque no lo hace de esta manera”. Entonces, yo genero un prompt técnico, me lo llevo al otro lado y le digo, “haga esto”, y hace exactamente eso y lo resuelve. Esa ingeniería de prompt es una herramienta fundamental para tener éxito cuando usted trabaja con herramientas agénticas.

Yo inventé una manera como para tratar de explicarlo, en donde eso que dice IAY, IAZ no es nada técnico, esto es un ejemplo que yo me inventé en donde IAY es el desarrollador y IAZ es un auditor y ese esquema que aparece ahí es cómo funcionaría en una posible aplicación en el SINAES. Utilizando ecosistemas de auto mejora o lo que se llama la ingeniería automática de prompts, el APE.

Lo importante es lo que dice ahí abajo, esto es fundamental, ese cuadrado al final le responde a una pregunta que a mí me hizo Geovanni, porque dice: Este paradigma dota a la ingeniería de prompts de la misma robustez y rigor técnico que la ingeniería de software clásica, minimizando el error humano. No sustituyéndolo, minimizándolo.

Aquí caemos en el tercer aspecto, dice arquitecturas RAG en entornos de alta privacidad, este es el tercer aspecto que es fundamental que tomemos en cuenta. ¿Qué es RAG? Retrieval augmented generation o lo que en español se traduce como generación aumentada

por recuperación. ¿Qué es eso? Básicamente, es que un usuario captura información que lo manda una base de datos y que utiliza lo que se llama un LLM como cerebro, ¿qué es un LLM? Bueno, es el modelo de gran lenguaje, ¿qué es un modelo de gran lenguaje? Los modelos de inteligencia artificial que se están desarrollando hoy en día. LLM es todos esos modelos que se han desarrollado como Anthropic, Gemini, todos los que se han desarrollado.

¿Qué es lo que pasa? Que esos modelos todo mundo cree que el que usted debe usar es el más potente y por eso es que son animales peligrosos y por eso es que los cuidadores también son peligrosos, porque están tratando de demostrar que su animal es el que más puede hacer y nos engañan muchas veces con eso. Pero resulta que hay muchos modelos que ya no son LLM, sino que son SLM Small Language Model, en donde yo puedo tomar uno de esos modelos, que es tremendamente potente, que ha sido desarrollado y que ha sido liberado y por eso se llama modelos outsourcing. Entonces yo puedo utilizar un modelo Llama, por ejemplo, que está a la disposición, puedo generar una estructura, que ya no está conectada con la nube, sino que la tengo on-premise y que yo desarrollo mi ecosistema, hago una curación de insumos. ¿Qué ventaja tiene y qué desventaja tiene? La ventaja que tiene es que ya no está contaminado con toda la basura que está ahí afuera, con todos los algoritmos que están sesgados, trabajo en un ambiente clínicamente más limpio.

Yo aquí lanzaría a esto como una, y por eso es que al final yo no digo que es algo exclusivo del SINAES, sino que es algo en lo que la academia debería de estar en este momento discutiendo mucho, pensando mucho, ¿por qué? Porque, si usted tiene muchos modelos de estos y después los interconecta, que eso sí es posible. Entonces usted estaría hablando de insumos muy protegidos, o sea, los modelos RAG no alucinan o alucinan muy poco. Los modelos RAG usted los modela para que se comporten de la manera más adecuada para sus propósitos, no son impermeables, usted los puede hacer permeables, pero con cuidado. Para que tengamos una claridad, ¿a qué se parece eso? A la herramienta NotebookLM, que antes se llamaba así y que ahora se llama Gemini Notebook. ¿Qué ventaja tiene esa herramienta? Esa es una herramienta RAG, solo que es una herramienta comercial y esa es la herramienta que yo usualmente uso más, porque es la que me permite sacar, extraer, hacer ese query de información y analizar solo eso. Yo me acuerdo que cuando se empezó a hablar de estos modelos de inteligencia artificial, lo primero que yo me imaginé es, bueno, qué interesante sería que yo pueda analizar solo ese contexto y que no se contamine con todo lo demás, porque yo le hacía preguntas y siempre me respondía en función de un montón de otras consideraciones que yo no le estaba preguntando, aquí es así. Entonces, esta variable de arquitectura RAG es fundamental.

El primer anexo, no sé si lo llegaron a ver, sé que es muy aburrido y es muy largo, pero cuando uno lo va leyendo, aunque no lo entienda mucho, cuando uno lo va leyendo es muy interesante porque se va dando cuenta de cómo usando todas esas variables técnicas usted lo va modelando, lo va construyendo específicamente para el SINAES, eso es lo que dice ese anexo. Entonces, usted va viendo cómo lo puede construir y ahí se hacen consideraciones de costo, por ejemplo, cuánto nos costaría tal cosa y salen opciones, lo hago on-premise o si lo hago en la nube, que es también se puede. Hay un gráfico que habla de a dónde se cruzan y dónde empieza a hacer mejor uno que el otro y así se va construyendo. Yo eso se lo enseñé a Geovanni, porque él me hizo la pregunta, ¿cómo hago yo para construir esto? Geovanni se imaginaba que una aplicación de estas era que yo le pongo prompts generales.

Dra. Susan Francis Salazar:

¿Falta mucho, don Ronald?

Dr. Ronald Álvarez González:

Yo puedo cortar en cualquier momento.

Dra. Susan Francis Salazar:

Es para que pueda presentar toda la presentación.

Dr. Ronald Álvarez González:

Ya más concreto, esta parte en donde hay una autoridad humana, por eso es humano-céntrico con todo este análisis colaborativo y por eso se habla de una gobernanza bio-cibernetica, porque es lo que se llama HAACS es una colaboración entre el diseño combinado de seres humanos con tecnología, eso es lo que significa. Más adelante habla de las famosas Redes de Petri, son ni más ni menos que notaciones, eso tal vez le haga sentido a Gerardo, que es ingeniero eléctrico. En donde era una anotación que se había inventado para poder escribir, aquí no lo puse, pero está en el documento, el esquema desde ese modelo HAACS con Redes de Petri, cómo queda al interno del SINAES, está en el documento. Aquí lo menciona, la aplicación del modelo de acreditación 2025 con Redes de Petri, con plazas de automatización y con transiciones con control humano.

Dra. Susan Francis Salazar:

Don Ronald, solamente una pregunta. Si voy leyendo bien, sobre todo por la presentación en donde aparecen los 5 elementos que usted expuso. Esa sería la línea para decir conductora de la ponencia, el equipo de trabajo tendría que darle ahora, para decir así, el contenido de la parte académica, ¿eso es lo que usted nos propone?

Dr. Ronald Álvarez González:

Sí, lo que está en el documento, en el documento tiene el respaldo de la parte académica de las consultas a las fuentes.

Dra. Susan Francis Salazar:

Ok. Más que la fundamentación de esto la parte ya de la propia dinámica de SINAES, porque es como es una ponencia de como SINAES se imagina la implementación de esto como una estructura ya formal, esa parte es lo que haría falta. Yo no tuve el chance, me disculpo, de leer ese documento. Estoy pensando que la ponencia, según lo que dice INQAAHE, debe ser accesible a una audiencia académica de diferentes contextos. Este elemento que nos está presentando usted de pronto es muy técnico y entonces habría que colocarlo como en un ambiente, voy a decirlo así, de las interacciones y de las acciones propiamente del SINAES para que haga sentido en esas audiencias. ¿Eso es lo que haría falta o el documento que ya usted nos propuso en la carpeta es el documento?

Dra. María Eugenia Venegas Renault:

Para mí es muy denso todo esto, toca entrarle a estudiar.

MAE. Marianela Núñez Piedra:

Susan, yo levante la mano. Cuando des el orden para hacer un par de consultas.

Dra. Susan Francis Salazar:

De acuerdo. Don Ronald, si a usted le parece, finalice porque yo lo interrumpí, con toda la pena, y para dar la palabra.

Dr. Ronald Álvarez González:

Esa matriz RACI, que es muy utilizada en estos modelos, tiene estos elementos que están a la izquierda, indexación y recuperación de evidencia, análisis de cumplimiento de las pautas, generación de dictamen preliminar, fallo final de acreditación. Es como en las etapas del SINAES actuarían tanto el agente de inteligencia artificial, el equipo técnico de la DEA, el par evaluador y el Consejo. Esto, Susan, sería un poco tratando de responder a su inquietud, es la idea de aterrizarlo con esta matriz en el modelo que se está buscando. Tratando de responder a la consulta que me hizo el equipo de TI ellos no veían cómo se implementaba, y tomando en cuenta lo que decía esta parte, que este paradigma dota a la ingeniería de prompts de la misma robustez y rigor técnico que la ingeniería de software clásico, minimiza el error humano, es como lograr utilizar la inteligencia artificial bajo todos estos principios en un ambiente cerrado, no es poner prompts ahí a lo loco, sino que ya nos va a estructurar un proceso para lograr cumplir con los objetivos específicos del SINAES. Aquí viene la parte de los desafíos epistemológicos del uso de esta tecnología, no me voy

a meter aquí, ustedes seguramente lo vieron. Al final, eso es como el gran objetivo, el inmenso potencial transformador de la inteligencia artificial agéntica en la educación superior solo se materializará si logramos subordinar la autonomía de la máquina a la reflexividad humana. No es solo que tome la decisión o que pueda evaluar, no, es a la reflexividad que esa gran incógnita de Carmody Grey. Y el SINAES no adoptará la inteligencia artificial para sustituir el juicio experto, sino para amplificarlo y creo que eso es bien importante.

Si ustedes me permiten en entrar en esta parte, aquí al puro final es lo que les quería mostrar, esto no está en el documento que yo les pasé, pero es el anexo tres y es el análisis de la propuesta desde la perspectiva de los insumos que nos pasó a la Universidad de Costa Rica de cómo queda esa evaluación. Haciéndolo muy rápido, dice, el análisis de alineamiento y cumplimiento, la coincidencia y áreas de alto alineamiento. Hay una preservación de la autonomía y la supervisión humana. Hay un enfoque estratégico de aumento de capacidades, capacity-up versus cost-out out. Hay una transferencia, explicabilidad y mitigación de alucinaciones, arquitectura RAG y analiza la propuesta nuestra y la propuesta del marco de referencia de la universidad. Hay una gradualidad y entornos controlados de prueba, sandbox, porque la propuesta dice que nosotros tenemos que ir llevando esto como planes piloto, no vamos a tirarnos de cabeza. El marco de resoluciones de la UCR coincide plenamente con eso. La protección de datos personales y la soberanía algorítmica.

En la propuesta nuestra promueve el despliegue de modelos de lenguaje pequeños de código abierto. El marco de resoluciones de la UCR cumple estrictamente con el principio de protección de datos personales, se apega a la ley, etcétera. Cada una de esas cosas está debidamente justificada.

También hay brechas entre uno y otro, que a mí me parece que es muy importante y que tenemos que tomar en cuenta. Una estructura de gobernanza e instancia de decisión. La Universidad de Costa Rica creó el Comité Estratégico Institucional de Inteligencia Artificial como órgano centralizado de dictámenes, seguimiento y aprobación de lineamiento. La propuesta del SINAES menciona a la matriz RACI, pero no formaliza una estructura orgánica interna. Y bueno, creo que eso es un buen elemento que tenemos que tomar en cuenta.

Otra brecha es el modelo institucional de control por etapas, este me parece muy interesante. El marco de la UCR estructura un ciclo de vida obligatorio para cualquier solución de IA, aquí lo planifica. Un pilar de alfabetización y formación continua. La propuesta nuestra sí lo dice para el Consejo y lo dice para toda la organización, pero no se pone como pilar de alfabetización y formación continua, sino la propuesta nuestra habla de capacitación en todos estos recursos. Me parece que sí lo cumple, yo no estoy tan de acuerdo que ahí tengamos una brecha. Luego mecanismos de impugnación y revisión humana directa, como esto fue analizado en un entorno RAG y aquí no toma en cuenta las mayores instancias del SINAES, en realidad sí las tenemos porque aquí habla de mecanismos de impugnación y revisión humana directa y nosotros tenemos el recurso de reconsideración. Entonces, aunque aquí lo ponga como una brecha, en realidad ya lo tenemos. Luego las dimensiones transversales, susceptibilidad, salud mental y género, esto me parece muy importante que lo tomemos en cuenta. La sustentabilidad, la UCR exige medir la huella de carbono y el consumo energético de la estructura y ejecuciones de inteligencia artificial. La salud mental, la UCR previene riesgos como el tecnoestrés, la fatiga digital y la pérdida del sentido comunitario mediante el enfoque “Una sola salud”. La perspectiva de género e inclusión, la UCR insta a realizar auditorías de sesgos de género y accesibilidad. Dice ahí que las anteriores tres dimensiones no son abordadas explícitamente el documento del SINAES, si eso es cierto.

Luego hay unas conclusiones, no las voy a leer, pero sí hay tres recomendaciones que me parecen, bueno, no tres, hay muchas, hay seis. Habla de crear un comité órgano de gobernanza de la IA en el SINAES, eso es algo lógico. Adoptar el modelo institucional de implementación y clasificación de riesgo, nosotros lo tenemos, pero específicamente para la IA. Formalizar el derecho de impugnación y auditoría externa, también lo tenemos. Estructurar un plan permanente de alfabetización en IA. Incorporar criterios transversales de sustentabilidad. Aprovechar los espacios de cooperación interinstitucional, y aquí dice, acoger la oferta formal de colaboración remitida en el oficio tal por la vicerrectoría de docencia para construir marcos de referencia y buenas prácticas. Hasta aquí sería.

Dra. Susan Francis Salazar:

Muchísimas gracias, don Ronald, sobre todo porque nos propone una estructura muy fuerte, yo diría que altamente desarrollada. Entonces, por eso la consulta que le hacía ahora al inicio.

Voy a dar la palabra entendiendo que ya estamos sobre las 12 del mediodía y que por lo tanto la propuesta para este punto sería establecer si adoptamos esta línea como la propuesta o si hacemos algunos ajustes para que podamos participar como grupo de trabajo algunas otras personas, yo los escucharía. De nuevo, les reitero, estamos sobre las 12:10 p.m., así que entonces las aportaciones ojalá sean lo suficientemente concretas para que tomemos la decisión. Tengo a doña Marianela, a doña Lady y a Melania. Adelante, doña Marianela.

MAE. Marianela Núñez Piedra:

Brevemente, yo probablemente por un error de lectura mío no entendí que esta presentación de don Ronald fuera, no sé, la guía, pensé que más bien nos estaba haciendo un poco su planteamiento acerca de lo que debe hacer SINAES, hasta ahora, después de la pregunta de Susan me confundí y entiendo que esta es su propuesta para presentar, valga la redundancia, nuestra propuesta para este congreso.

Yo entiendo la perspectiva de don Ronald, pero la veo desde un punto de vista teórico, máxime cuando él mismo dice que el equipo de TI de SINAES tiene serias dudas respecto a cómo implementar la inteligencia artificial y la transformación tecnológica a lo interno de la institución.

En mi experiencia en este tipo de congresos, tal como lo dice el call for proposals que nos presentó INQAAHE, ellos dicen que el objetivo del evento es capacitar a las instituciones de educación superior para responder eficazmente a un entorno tecnológico, es decir, toda esta propuesta de don Ronald efectivamente cumple con esto y está bajo la línea de los cuatro ejes temáticos que ellos plantean. Mi pregunta es si nosotros estaríamos pensando en ir a presentar un deber ser y no lo que tenemos, cómo ha funcionado y los espacios de mejora, porque lo sabroso de este tipo de eventos es ir a compartir las mejores prácticas, no lo que yo creo que deberíamos tener, pero no tengo. Ahí está mi confusión y es lo que quería nada más mencionar.

Dra. Susan Francis Salazar:

Gracias, doña Marianela. Adelante, doña Lady.

Dra. Lady Meléndez Rodríguez:

Lo mío es más puntual en relación con el contenido propiamente por una experiencia que escuchamos de la Universidad del Valle de Guatemala, donde ellos hacen una aplicación muy cercana a esta, pero precisamente para procesos de aprendizaje. Ellos hablaban de que los agentes en la universidad tienen el problema cuando se trata de datos perimetrados, de que se saturan y entonces se vuelven incapaces, por ejemplo, de actuar sobre situaciones emergentes, sobre imprevistos, sobre excepciones y para innovar. En ese caso pareciera que estamos hablando de agentes versus human in the loop, porque más bien parecen alejarse de esa propuesta más humanizadora de lo que estaríamos planteando.

Lo otro es que también se hablaba en otros contextos del riesgo, de más bien deshumanización epistemológica. A lo que se refieren es que cuando estamos hablando de un uso estocástico de la inteligencia artificial, donde ya se predican riesgos, por ejemplo, en la evaluación de los aprendizajes, porque los estudiantes hacen la misma pregunta y es la IA la que responde y ellos simplemente reproducen la misma respuesta que les da la IA, que cuando se hace un trabajo de simulación epistemológica donde hay una mayor generación de conocimiento, más bien es todavía más grave porque se vuelve sustitutivo de la propia persona de qué va a ser, además después con la respuesta que le está dando la IA generativa, porque ya ni siquiera es que ellos hacen algo con la respuesta que les da, sino que la misma IA es la que sigue haciendo cosas por sí misma con esa respuesta. Entonces, entre más sofisticada la inteligencia artificial, más sustitutiva del conocimiento humano se vuelve. Entonces, que intentando ser hacer un uso más interactivo, generativo en los contextos de la educación, más bien se está corriendo el riesgo de que cada vez es más sustitutivo del propio desarrollo del conocimiento de los mismos estudiantes. Entonces, ¿cómo neutralizar ese tipo de cosas? No para que lo responda, por supuesto, don Ronald, sino contándole cuáles han sido algunos de los principales cuestionamientos que están surgiendo en estos momentos de los contextos educativos, por si tenemos alguna respuesta que dar en ese contenido. Gracias.

Dra. Susan Francis Salazar:

Gracias, doña Lady. Melania.

PhD. Melania Brenes Monge:

Yo en principio agradecerle a don Ronald, a mí me parece que esta propuesta y este análisis que él nos presenta es oro puro. Yo no creo que haya otras agencias o sistemas pensando con este nivel de especificidad, creo que es bastante adelantado creo que la discusión que se está dando en este momento en donde todavía están como enclochados en el marco de si vale la pena, si avanzamos. Y uno lo lee desde incluso en documentos internacionales de entidades que deberían ya estar aterrizando las cosas en este nivel.

A mí, me parece que por lo que yo le entiendo y creo que viendo también como una especie de outline para una ponencia, yo creo que aquí don Ronald lo que nos presenta es todo un enfoque de aprovechamiento educativo de la IA desde la visión de la evaluación y acreditación de la calidad, centrado en el humano, utilizando herramientas propias, generando un rol importante a quien analiza y no el tema de este de capacity in. Todo ese enfoque de principios de aprovechamiento dentro del marco de nuestro negocio, llamémosle así, que es la evaluación y la acreditación y la investigación y la generación de conocimiento, me parece que es un marco específico de principios que podrían orientar muy claramente una política de cómo utilizar la IA internamente. Entonces, a mí por lo menos me resulta muy útil como para revisar, un insumo más, ya tengo muchas observaciones, un insumo más para lo que habíamos pensado que podía ser esa ruta de transformación digital, después vino lo de la Contraloría, después vino esto, pero son muchas cosas que pueden nutrir eso pero mirándolo como propuesta de ponencia, yo les propondría que pudiésemos entresacar comunicativamente esos principios para inclusive someterlos a consideración para ver qué están pensando otros sobre esto, si están pensando parecido o están pensando diferente y un poco también incluso puede ser emparejarlo a la transformación hacia un nuevo modelo que vamos a transitar el año entrante, porque un poco aquí la complejidad de lo que yo veo que nos está proponiendo don Ronald en términos de usar agentes, es que esto tiene gente que lo entrena, es decir, somos nosotros mismos, no es un otro que nos ofrece un servicio, nosotros dedicamos tiempo a educarlo y de la calidad de los datos que se incluya ya sea que alucina. Creo que un muy buen inicio con la nueva información que viene del modelo 2025 puede ser incluso algo más curado en términos del dato que lo que teníamos sinceramente del modelo que estamos abandonando este año, que sí podría generarnos algún tipo de problema, porque

si bien, como dice doña Marianela, en estos eventos lo ideal es que uno pudiese obviamente llevar como una puesta en práctica y hablar de los logros y los desafíos.

Sí, me parece que a este punto que estamos también es importante llevar nuestros marcos de cómo estamos pensando el aprovechamiento de esto institucionalmente en el escenario de una transición hacia un nuevo modelo que queremos que tenga otras características epistemológicas y que se pueda someter a consideración para que otros nos digan, “yo también voy por ahí” o “nosotros no vamos por ahí, apenas vamos pensando” o “nosotros nos vamos a orientar por un animal peligroso, que es tal herramienta que ahí es donde la vamos a usar”. Todo ese tipo de temas creo que podríamos abrirlo como más un enfoque de diálogo y de realimentación.

Dra. Susan Francis Salazar:

Muchísimas gracias, Melania. Yo vería dos caminos, desde la invitación que se hizo en este Consejo no hay otra propuesta de estructura de ponencia. Recordar nada más que la fecha límite que se da desde INQAAHE es al 31 de octubre, que lo que habíamos acordado era crear una posición institucional, que fuese un grupo de trabajo institucional que trabajara un tema y que ese tema se propusiera como ponencia. Si no existiese otra propuesta de parte de los miembros de este Consejo, yo abogaré por aprobar o que estuviéramos de acuerdo con la propuesta que hace don Ronald.

Yo sí creería que en términos de comunicación necesitamos bajarle el nivel técnico y enfocarnos más desde la naturaleza propia de la función que tiene SINAES y que sean los principios los que nos permitan ir concretando cómo lo está pensando SINAES y cómo lo concretaría y ahí sí se pueden ir entremezclando las cuestiones técnicas. De hecho, a mí me encantaría participar de ese grupo de trabajo para poder entonces ahora sí materializar esa función sustantiva que tiene SINAES con este tipo de discurso técnico. Eso es un poco, sobre todo pensando que la invitación que hace o el llamado que hace INQAAHE lo subraya, que la ponencia tiene que ser como muy accesible.

Trataría además de dar respuesta a lo que ya doña Lady señaló, de cómo resolver ese tipo de cuestiones, que según lo que entendía aquel día, que también lo escuchamos ambas, me parece que tiene que ver mucho con la formación de capacidades de la organización, más allá de cómo se programa. Yo voy a hacerle muy franca, me gusta mucho el enfoque, porque si algo escuché ese mismo día y quedé muy clara, es que los enfoques de uso de inteligencia artificial en cualquier organización en este momento se están pensando desde la potenciación de las capacidades, cómo llegar a ser más productivo, cómo poder innovar, cómo ser creativos más allá de cómo resolver tareas. Entonces, me parece que el enfoque lo da, solo que yo, con todo respeto, le bajaría un toquecito ese tono así tan técnico y la mirada que yo le propondría es que lo miremos desde la acción sustantiva del SINAES, porque todos los títulos que usted coloca sí dan para poder hacer esa armonización. Don Walter, dígame usted.

Ing. Walter Bolaños Quesada:

Muchas gracias. No hay que olvidar que ellos van a analizar también la propuesta, ellos van a ir viendo si calza o no con los intereses que tienen propuestos.

A mí, particularmente me gusta la opción, creo que sí es importante que no sea tan técnica excepto que usted pueda o vaya a tener suficiente tiempo para ir explicando un poquito más esos detalles que en general la gente no conoce y entonces eso la aterriza mucho más para comprenderla mejor de parte de todo el auditorio.

No hay que olvidar tampoco que ellos están pensando en lo que decía un poco doña Marianela en buenas prácticas o de alguna manera que se vaya a exponer una experiencia, una investigación ya vivida, que ya tiene resultados. Sin embargo, por ser este tema, yo estoy de acuerdo con lo que dice Melania, en realidad todavía es muy posible que no haya experiencias suficientes para exponer y en temas como estos, con la velocidad que avanza, difícilmente uno está al día, usted la prepara hoy y ya mañana es tarde para exponerla.

Entonces, sí es bueno que la vean de ese modo y va a ser una apuesta que se vaya a tener en el SINAES y que después sea o haya otra oportunidad para explicar los resultados que pueden ser de mucho beneficio para las otras agencias, eso sería interesante para el SINAES mismo, porque ir a la cabeza en esto no es sencillo y yo creo que el trabajo que ha hecho Ronald es valioso y vale la pena que si ahí le dan la oportunidad la exponga y después se puedan presentar ya los resultados finales que se tengan de la puesta en marcha de esto. Gracias.

Dra. Susan Francis Salazar:

Muchas gracias, don Walter. De eso que usted está planteando, yo quisiera retomar que la invitación o el llamado dice que estamos buscando propuestas críticas y orientadas por la práctica que permitan repensar el aseguramiento de la calidad como un proceso centrado en lo humano dirigido hacia el impacto y hacia el futuro. Y después dice, las propuestas pueden ser basadas en nueva investigación, prácticas innovadoras u otras experiencias relevantes, experiencias e iniciativas. Y yo vería esto como una iniciativa, me parece que sí coincide o sí da. Además, para plantearle a doña Marianela, las ponencias tienen un alto contenido de propuesta, entonces yo vería que esto es como un marco como para pensarse la propuesta en términos de todo lo que implica una iniciativa cuando se tiene que fundamentar. Doña Lady.

Dra. Lady Meléndez Rodríguez:

Una idea que me parece que podría ser es una propuesta matricial de procesos que ya existen o que existen sobre la realidad y cómo se verían transformados desde esta aplicación. Y lo otro es hacer un ejemplo en vivo.

Dra. Susan Francis Salazar:

Es lo que don Ronald fue colocando.

Dra. Lady Meléndez Rodríguez:

Sí, pero ejemplificado con un procedimiento real para que se vea cómo es que funciona.

Dra. Susan Francis Salazar:

Sí, tal vez sacado de la lámina técnica. Don Gerardo.

M.Sc. Gerardo Mirabelli Biamonte:

Me parece que es un tema muy interesante, es un tema que va a ser bien recibido porque me parece que está bien desarrollado. Yo creo que, dentro de ese grupo, de esa Comisión, se debe analizar primero qué es lo que queremos hacer, porque aquí se habla de que puede ser una sesión de hasta 55 minutos o puede ser un póster, definir por dónde le queremos entrar, porque si es una sesión ya si hay una responsabilidad mayor. Debería plantearse definitivamente como una iniciativa, porque es algo que nosotros todavía estamos discutiendo y debe quedar muy claro de que, si bien es cierto, hay todo un análisis, ese es un análisis que debe ser todavía comenzado a implementar de alguna forma dentro del SINAES. Y también, debe trabajarse en hacer la presentación de la propuesta, porque ellos piden una serie de elementos antes de poder determinar si lo aceptan o no lo aceptan, creo que es parte del trabajo de la Comisión. Y yo también me apunto en la comisión para poder conversar.

Dra. Susan Francis Salazar:

Antes de que don Ronald haga el cierre, quisiera ofrecerme, sobre todo para precisar algunas cuestiones que yo he estado tratando de, pero me parece que usted ya las tiene muy bien dibujadas. Creo que Kattia, sobre todo por los últimos estudios que le escuché, está haciendo, sería una participante muy interesante en esto. Le escuché haciendo una reflexión sobre marcos de gobernanza para inteligencia artificial para la institución y me llamaron poderosamente la atención. Yo creo que definitivamente tendríamos que pensar en alguien de la DEA, no sé, doña Melania, si tiene algún tipo de perfil de una persona que pueda tener sensibilidad al tema y que nos pueda aportar. Y yo creo, no sé hasta dónde,

pero me parece que también para no hacer el grupo muy grande, deberíamos pensar en alguien de tecnologías, a no ser que don Ronald diga que no.

La Dra. María Eugenia Venegas Renauld, se retira a las 12:30 p.m.

MAE. Marianela Núñez Piedra:

Yo sigo pensando que la propuesta debe tener alguna vinculación o por lo menos alguna relación con la parte práctica. Quizás quien más utilidad podría darle a esta propuesta de don Ronald es alguien de la DEA, creo que ese sería el perfil idóneo como para poder aterrizar y al menos poder decir y estamos trabajando estas propuestas en x,y, pero me parece que la DEA es el que más se beneficiaría.

Dra. Susan Francis Salazar:

Exacto, doña Marianela. Eso justamente estamos diciendo.

Dr. Álvaro Mora Espinoza:

Yo lo que propondría es que se le indique a la dirección de la DEA que puede designar alguna persona de esa instancia y sus funcionarios también se sientan incorporados.

Dra. Susan Francis Salazar:

Sí, de acuerdo. Adelante, don Ronald para que usted cierre la sesión.

Dr. Ronald Álvarez González:

Yo creo que todos los comentarios son lógicos y tienen razón de ser, dudas, desde la perspectiva de que si esto puede ser una buena idea para una ponencia. Yo les confieso que para mí eso fue secundario, en realidad, para mí eso no fue como lo más importante, mi objetivo no era eso. Yo no sé ni siquiera si voy a estar aquí el año entrante, porque creo que esto no va a ser este año, sino el año entrante. Yo lo que me propuse fue responder a cómo el SINAES debe de encontrar una forma que tome en cuenta todos esos elementos que se han detectado ahí como fundamentales para poder implementar esa tecnología de una manera adecuada, para mí eso era lo importante, para responder a la pregunta fundamental que a mí me hizo Geovanni, porque cuando yo hablé con Geovanni ya yo había avanzado bastante con esto. Entonces, ahí fue donde yo me imaginé ese modelo de aplicación que es el anexo 1, porque ahí en las respuestas técnicas, porque usted no puede hacer esto sin apartarse de la parte técnica, porque si usted no lo puede demostrar desde el punto de vista técnico, cualquiera le dice, “sí, pero ahí no está considerado el modelo SLM que usted va a utilizar, ahí no está consideradas las tarjetas de procesamiento que se requieren, ahí no está considerado la capacidad de almacenaje que se requiere, ahí no está considerado una serie de aspectos técnicos”. Entonces, en el anexo se aclara, este es un ejemplo de cómo sería una implementación específicamente para el SINAES. No puede jamás ser considerada como la base para hacer la implementación, aquí requerimos el personal técnico que domine adecuadamente estos temas y todo su planteamiento para la implementación tiene que ser revisado.

Desde la posición en la que estamos ahora, con un proyecto de transformación digital en donde ni siquiera ha sido aprobado por el Consejo, estamos hablando de una tecnología que no ha terminado de configurarse, pero que es por todas las consideraciones que se hacen a nivel mundial, una tecnología que se las trae, que nos va de alguna manera a impactar en el futuro. Y que, si nosotros no hacemos algo, lo peor es no hacer nada, podemos no hacer nada, pero es la peor decisión que podemos tomar. Entonces, yo me vengo rascando la cabeza pensando, ¿cómo armamos este muñeco tan complicado? Y sí, todas las preguntas que ustedes se hacen surgen.

A mí, se me ocurrió que esta podría ser una ponencia, sí, porque se planteó al interno del Consejo y la primera observación de Marianela, que es absolutamente lógica, no tiene, como lo mencionó Melania, una posible aplicación, porque esto va tan rápido, por lo menos en la investigación preliminar que yo hice no aparece nadie, ninguna agencia de acreditación, y si existiera en alguna parte del mundo lo hubiera dicho. Hay aplicaciones en

agencias de acreditación que parcialmente lo utilizan. Hay aplicaciones en instituciones universitarias que también lo utilizan.

Me gustó mucho estas preguntas de doña Lady, porque resulta que depende del tipo de organización sus retos son diferentes. Los retos que se está hablando aquí de los agentes versus la saturación se tocan en el informe técnico este que yo construí como uno de los elementos que tiene que ser considerado y tiene mucho que ver con hasta dónde usted permite que los agentes actúen como entes independientes y hasta dónde llegan para que el proceso de enseñanza sea continuado por el docente, por ejemplo, que eso lo cubre este marco de control de la UCR, ahí se toca.

El Ing. Walter Bolaños Quesada, se retira a las 12:37 p.m.

La simulación epistemológica es fundamental también hasta cierto punto, o sea, usted no puede soltar a los agentes a lo loco porque los agentes padecen del síndrome de lo que se llama el principal agente y se vuelven independientes, y resulta que se olvidan de los prompts que usted le dio y empiezan a darle valor a la autoridad y resulta que hay un agente que es el que controla y ya no escuchan al humano, sino que escuchan a la autoridad superior y se rebelan, y hay unos experimentos terribles alrededor de eso, pero a eso hay formas de controlarlo.

Los retos que una agencia de acreditación como nosotros tiene son un poco distintos, es hasta dónde llega ese modelo HAACS red de Petri, hasta dónde lo configuramos para que sea híbrido, en el sentido de que hasta todos estos procesos de los que adolecemos y el tema de los insumos para poder nosotros aprobar a los equipos de pares evaluadores es un buen ejemplo. Usted con un buen esquema como este que se plantea en este modelo, usted puede encapsular ese proceso para que todos esos análisis engorrosos de comparar las competencias que debe tener los candidatos versus los planes de estudio versus todos esos diferentes tipos de circunstancias específicas de los procesos, usted esto lo puede estructurar de tal manera que le ayude al gestor para tener unos insumos más adecuados, pero al final le llega al Consejo un insumo y el Consejo analiza ese insumo, puede hacer que le que le quite, etcétera, y toma la decisión.

Dra. Susan Francis Salazar:

Don Ronald, vamos a ir terminando aquí dado por la hora. Agradecerle profundamente por la propuesta.

Dr. Ronald Álvarez González:

Perdón, solamente decir que la salvaguarda de los datos es uno de los elementos fundamentales, está considerado toda la ciberseguridad y el aspecto de confidencialidad está considerado en el modelo. Es fundamental porque la ley nos lo exige, entonces está considerado.

Dra. Susan Francis Salazar:

Muy bien. Entonces, si le parece, don Ronald, Marchessi y yo, una vez que doña Melania nos haya indicado a quién estaría sugiriendo la DEA, porque tendríamos para trabajar este mes de setiembre y tal vez parte de octubre para dar una versión y tomar esas decisiones que dijo don Gerardo, si vamos a optar por ponencia, vamos a optar por póster, todo ese tipo de condiciones.

Dr. Ronald Álvarez González:

Un último detalle, mencionó doña Lady la posibilidad de hacer una simulación y resulta que yo he hecho ya muchas simulaciones, entonces les puedo mostrar mucho de eso.

Dra. Susan Francis Salazar:

Me parece una muy buena idea hacer esas matrices, son bastante prácticas para didácticas. Con eso estaríamos cerrando la sesión. Muchísimas gracias.

Voto disidente en el artículo 3 del acta: Votación de declaración confidencial.

Dra. Lady Meléndez Rodríguez

SE CIERRA LA SESIÓN A LAS DOCE HORAS Y TREINTA Y NUEVE MINUTOS DE LA TARDE.

Dra. Susan Francis Salazar
Presidente

Mag. Marchessi Bogantes Fallas
Coordinadora de la Secretaría del Consejo