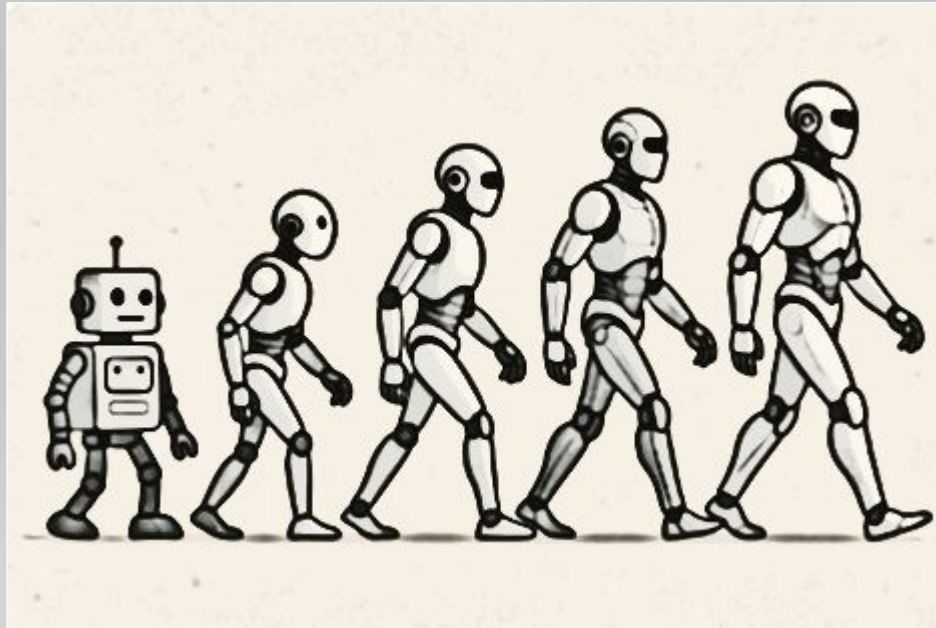


The Evolution to Agentic AI in Cybersecurity -

Threats, Opportunities, Responsible Use,
and the Race to Protect



Presented By: Tina Lampe



Tina Lampe
Germantown, Illinois

- Director of IT Software Engineering at DIRECTV
- MS Degree in Cybersecurity from Maryville University
- Assistant Vice President of the St. Louis ISACA board
- ISACA SheLeadsTech liaison
- Board member of St. Louis Infragard
- AAISM (Advanced in AI Security Management)
- CISSP, CISM, PMP
- USA Ambassador for the Global Council for Responsible AI

Please feel free to ask questions about any of these.

Points on our Journey:

1. AI Terminology and Output Testing
2. The Evolution: Harnessing the Power of AI in Cybersecurity
3. Adversary Amplification with AI
4. Ethical and Privacy Considerations
5. The Near-Term Future
6. Key Actions

1. AI Terminology and Output Testing:

- * **Generative AI**
- * **Large Language Models (LLMs)**
- * **Agentic AI**

GenAI, LLMs , Agentic AI – What's the Difference?

Gen AI	LLMs	Agentic AI
Can <u>Generate</u> new content such as text, images, music, code, video	<u>Large Language Models</u> (LLMs) are advanced machine learning models designed to understand and generate human language .	Also called <u>Autonomous AI</u> or <u>Self-directed AI</u> – is proactive without constant human guidance
<u>Reactive</u> – waits for input before creating something new	Humans interact with Large Language Models (LLMs) using <u>Prompts</u> - LLMs will process and generate <u>language or ideas</u> based on these prompts	Operates autonomously (makes decisions, executes actions, perceives then adapts to environment changes, processes information, generates outputs) <u>to achieve specific goals</u>
<u>Test Output Quality</u> - by validating output based on input request	Techniques to interact may involve <u>prompt engineering</u> and <u>RAG</u> (retrieval-augmented generation)	<u>Test Output Quality</u> – by focusing on objectives, constraints and oversight mechanisms.
	<u>Test Output Quality</u> - by validating output based on input request	

AI Agents –vs- Agentic AI

AI Agents	Agentic AI
❑ Example: Chatbots for customer support, simple anomaly detection systems ❑ Specialized AI programs designed for specific tasks ❑ Operate within predefined parameters ❑ React to inputs based on programmed rules ❑ Limited autonomy and decision-making capability	❑ Example: Autonomous threat hunting systems, adaptive defense mechanisms ❑ Advanced AI systems with a degree of autonomy ❑ Can make independent decisions and take actions ❑ Learns and adapts to new situations ❑ Operates across multiple domains and tasks ❑ Collaborates with other AI systems and humans
Source: https://right-hand.ai/blog/agentic-ai-in-cybersecurity/	



2. The Evolution:

- * **Harnessing the Power of AI in Cybersecurity**

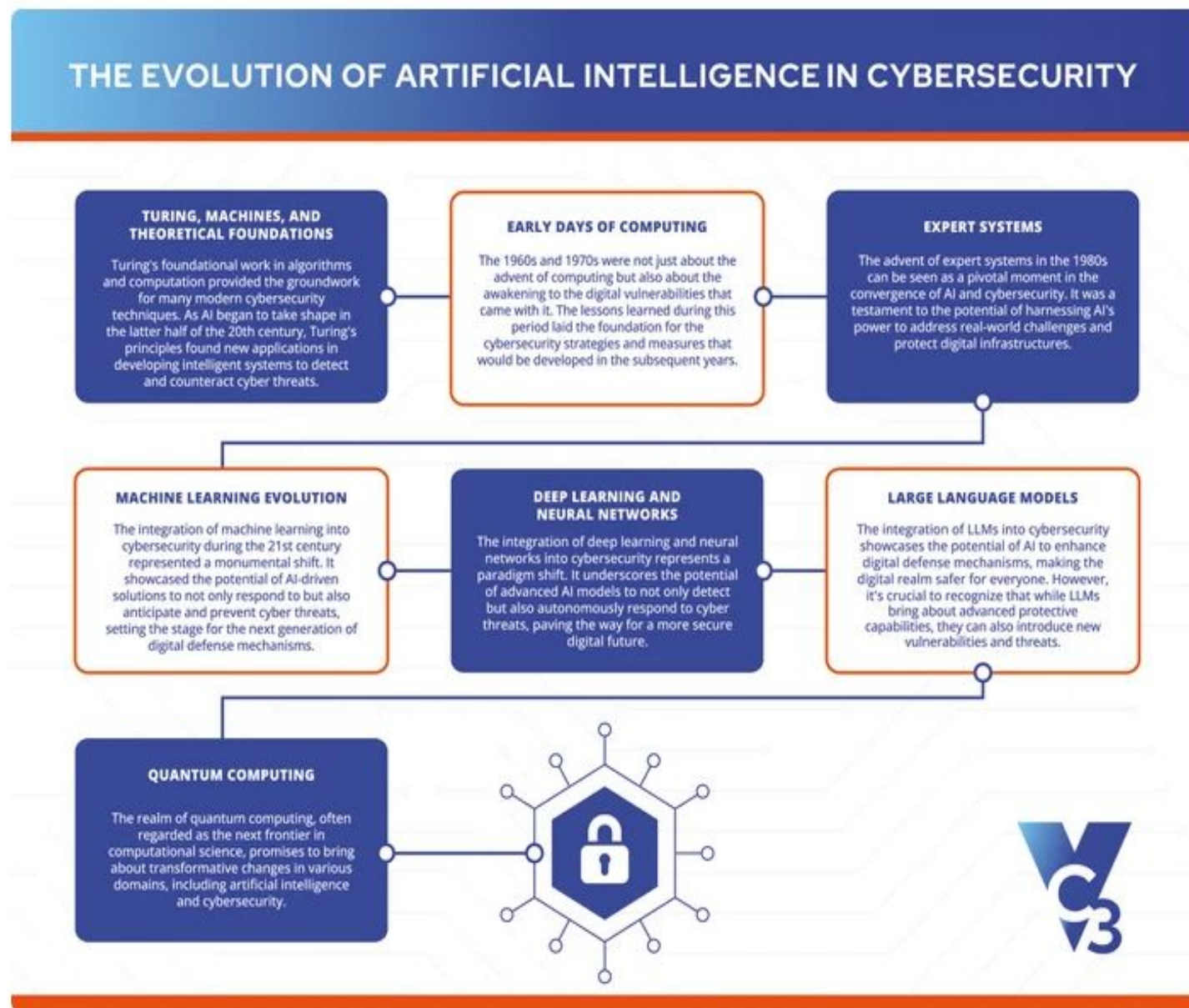
The AI/Cybersecurity Evolution

Approx 1980's:

- First computer Virus
- Cybersecurity as a profession

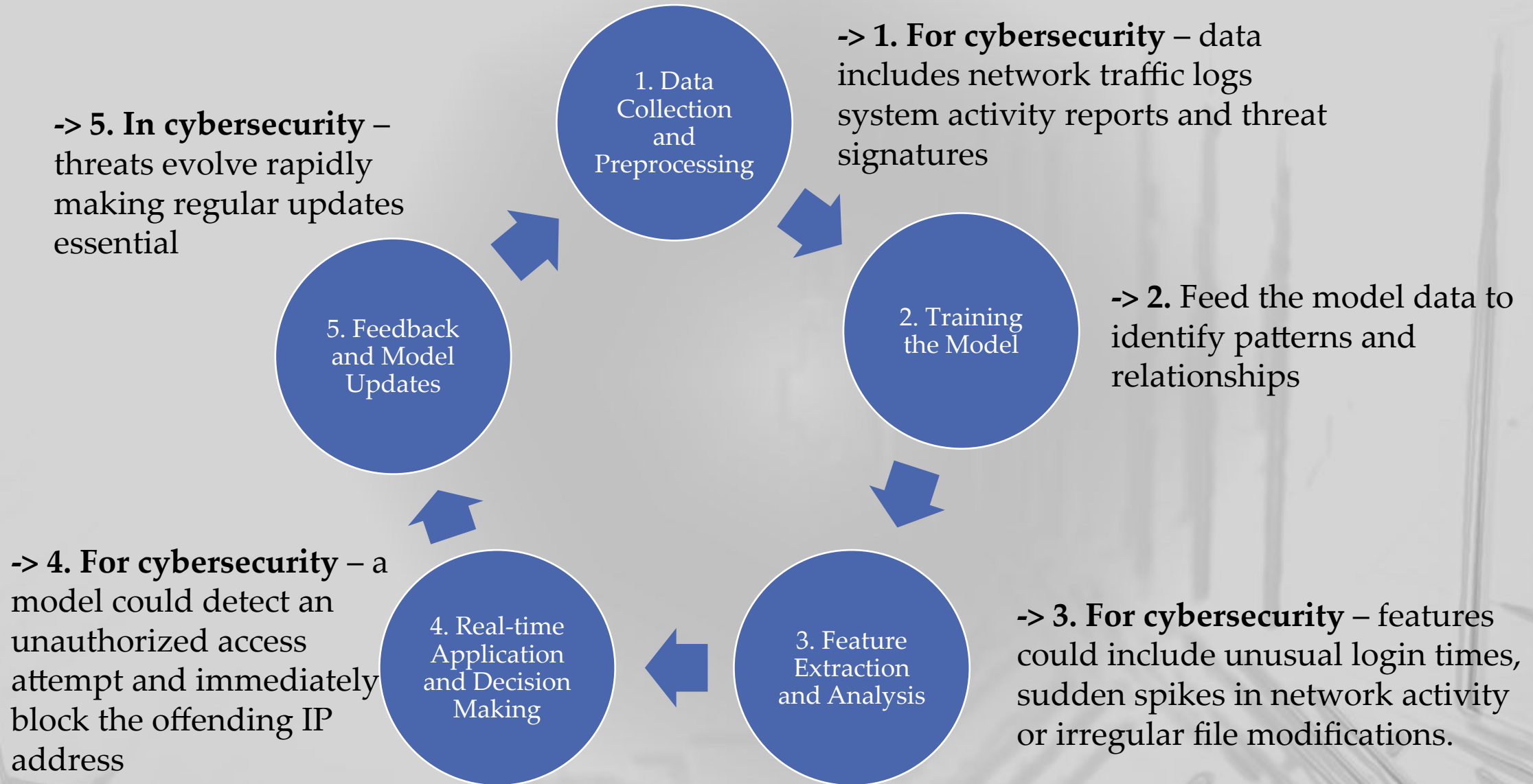
Early 2000's:

- Integration of AI via machine learning into cybersecurity



* Source: [vc3 blog](#)

Historic Machine Learning/AI Process



2024 View - AI Enhanced Defense

An Efficient, AI Enhanced SOC
'Finding the needle in the haystack'

Uncover hidden priority threats

Predict future attacks

Real Time Incident Response

Auto-Adapt defense tactics to evolving threats

Behavior Analysis

Zero Trust Architecture Support

Enhanced Human Oversight and Decisions

Intelligent Alert Prioritization



AI Powered Threat Intelligence Platforms
(aggregate and analyze data from various sources)

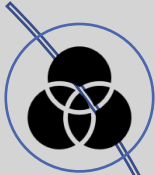
* Source: analyticsinsight.net

Continuous Improvement: Cybersecurity Machine Learning/AI Use Cases

CyberSecurity Machine Learning Use Cases

- Detect and Classify Threats
- Detect Anomalies
- Detect and Prevent Malware
- Intrusion Detection
- **Detect Spam and Phishing**
- **Endpoint Security**
- Network Risk Scoring
- Managing Vulnerabilities
- **DDoS and Botnet Protection**

More Recently – Cybersecurity Enhanced with GenAI



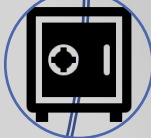
Data Masking for Privacy Protection



Synthetic Malware Generation



Reporting enhancements



Secure Code recommendations



Automated Secure Policy Generation



Creation of Interactive Cybersecurity Training

2025 – Towards more Agentic AI in Cybersecurity

Security Operations Agentic Use Cases	Application Security Agentic Use Cases
Triage and Investigation	Risk Identification
Adaptive threat hunting	Application test creation and adaptation
Response Actions	Dynamic application test execution
Adversary Simulation	Autonomous Test Case remediation
Remediate External Exposure	Automated Pen Testing

Additional Agentic Use Case Execution in CyberSecurity

SecOps	AppSec
□ Triage and Investigation □ Alert Deduplication □ Alerts Grouping □ Alert Enrichment □ Adaptive Threat Hunting □ Decomposing the Alert □ Searching for Atomic and Computed Indicators □ Analyze Behavior Indicators □ Response Actions □ Generating Infrastructure as Code □ Perform Endpoint Actions □ Security Controls	□ Risk Identification □ External and Internal Discovery □ Application Test Creation and Adaptation □ Adapting to UI changes □ Dynamic Application Test Execution □ Multiple Browsers, Devices □ Autonomous Reporting and Predictive Suggestions □ Autonomous Remediation and Test Case Correction □ Automated Pentesting

Source: <https://research.aimultiple.com/agentic-ai-cybersecurity/>

Challenges - Agentic AI

Agentic AI Cybersecurity Implementation Challenges

Lack of transparency and interpretability

Data Quality and Breadth concerns

Maintaining Reliability

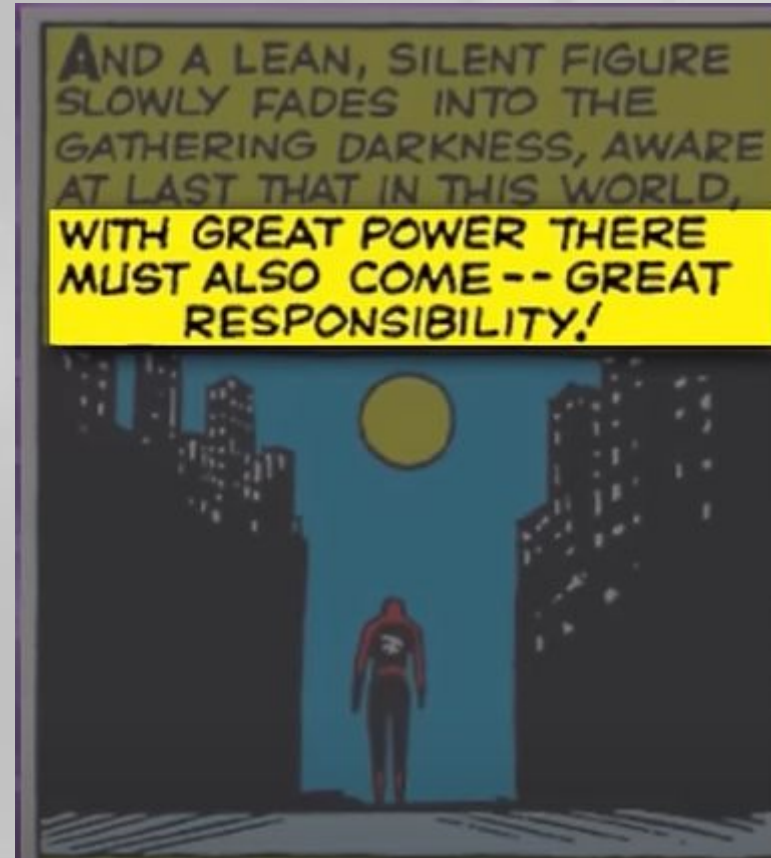
Complexity of Implementation

Human Oversight needs

3. Adversary Amplification with AI

Spider-Verse Wisdom

With great power comes great responsibility, and the ethical use of AI is paramount in securing our digital future.



AI Lowers the Barrier to Entry for Sophisticated Cyber Attacks

According to cybersecuritynews.com, “AI-assisted approach significantly lowers the barrier to entry for malicious actors and enables rapid scaling of attacks-even by technically unskilled individuals.”

* Source: cybersecuritynews.com

Major AI-Driven Threat Vectors





4. Ethical and Privacy Considerations

Ethical Considerations



- Privacy –vs- Security
- Bias and Fairness
- Accountability and Decision Making
- Transparency and Explanation

* Source: [isc2.org](https://www.isc2.org)

America's AI Action Plan

July 2025

“... To win the AI race, the U.S. must lead in innovation, infrastructure, and global partnerships. At the same time, we must center American workers and avoid Orwellian uses of AI. This Action Plan provides a roadmap for doing that,” said AI and Crypto Czar David Sacks.

* Source: www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf

... And the Race to Protect Our Organizations

- “AI systems are becoming essential tools, profoundly shaping how Americans consume information, but these tools must also be **trustworthy**.”
- Action Plan includes directives such as:
 - Bolster Critical Infrastructure Cybersecurity
 - Promote Secure-By-Design AI Technologies and Applications

* Source: www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf

AI Value Enablers – World Economic Forum

Cybersecurity is an integral part:

- Frameworks and guidance
- Audits and assessments

#Global Council for Responsible AI
US Ambassador

The AI strategy of an organization drives value alignment and must be evaluated through the AI life cycle. This can be achieved through a strong AI governance process. According to the World Economic Forum, there are various enablers to AI value alignment, including:

- **Frameworks and guidance** - foundational elements to guide AI solution development, deployment, and management
- **Human engagement** - essential to the iterative refinement and enhancement of AI solutions
- **Organizational change** - essential strategy to ensure alignment with organizational culture
- **Audits and assessments** - necessary to ensure the effectiveness of AI solutions in maintaining adherence to ethical values and best practices

Source: Larson, B.; Dignum, V.; "AI Value Alignment: How Can We Align Artificial Intelligence with Human Values," World Economic Forum, 17 October 2024, <https://www.weforum.org/stories/2024/10/ai-value-alignment-how-we-can-align-artificial-intelligence-with-human-values/>

Source: World Economic Forum, "AI Value Alignment: Guiding Artificial Intelligence Towards Shared Human Goals," October 2024, https://www3.weforum.org/docs/WEF_AI_Value_Alignment_2024.pdf

Trustworthy AI Framework Examples

NIST AI 100-1:



Fig. 4. Characteristics of trustworthy AI systems. Valid & Reliable is a necessary condition of trustworthiness and is shown as the base for other trustworthiness characteristics. Accountable & Transparent is shown as a vertical box because it relates to all other characteristics.

What are the main benefits of implementing ISO/IEC 42001? ^

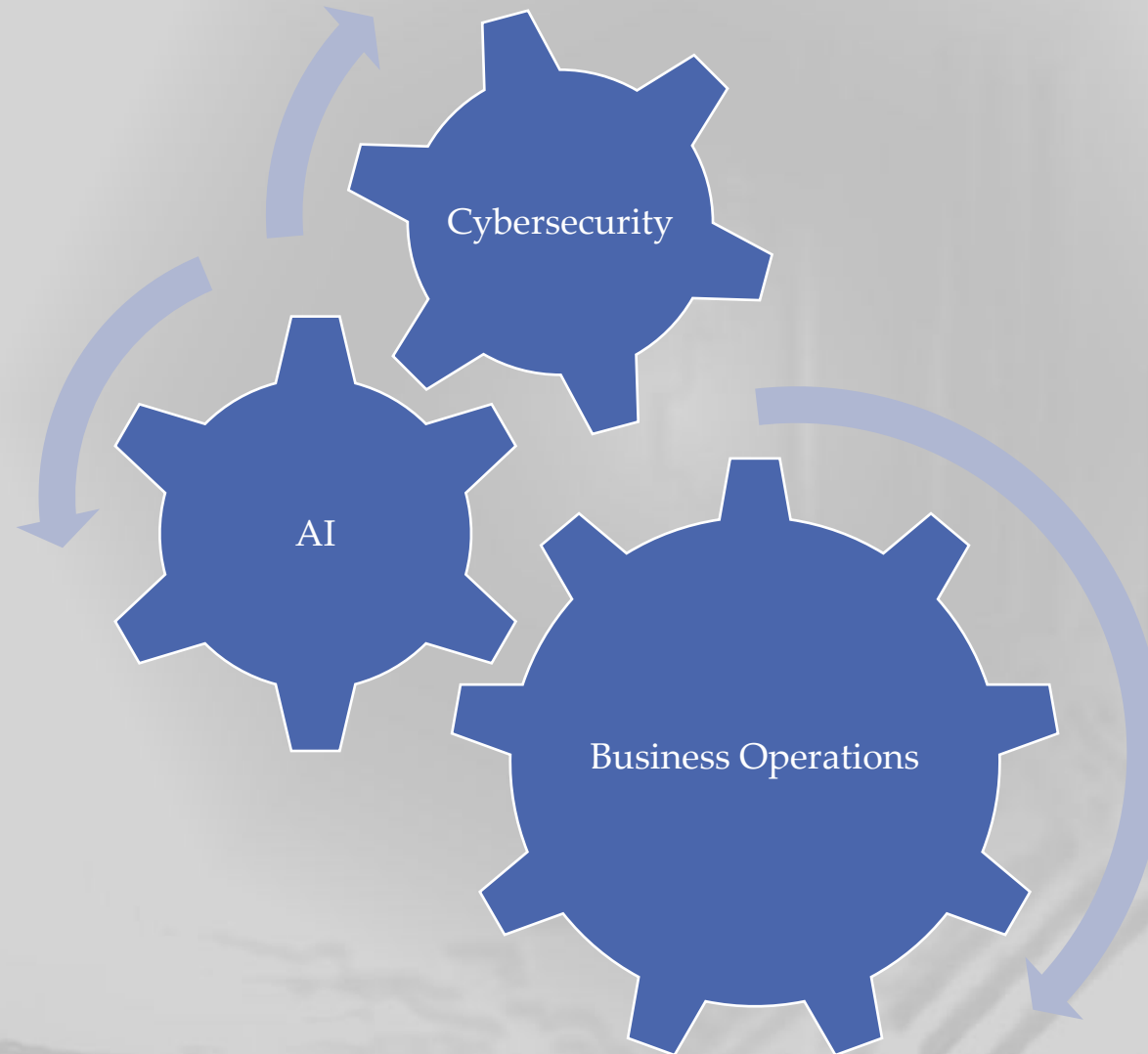
- **Responsible AI:** ensures ethical and responsible use of artificial intelligence.
- **Reputation management:** enhances trust in AI applications.
- **AI governance:** supports compliance with legal and regulatory standards.
- **Practical guidance:** manages AI-specific risks effectively.
- **Identifying opportunities:** Encourages innovation within a structured framework.

ISO/IEC 42001



5. The Near-Term Future

AI Becoming more Deeply Integrated into Business Operations



Continued Focus AI in Cybersecurity

Cyber Resilience and Agility

Targeted Organization
Rankings and Remediations
Based on Business Impact

Refinement of Cyber
Defenses near real time

Freeing the Security Team
to work on Strategic Tasks



Scaling Cyber Defense
Activities when Needed

Continuous Compliance
Assessments

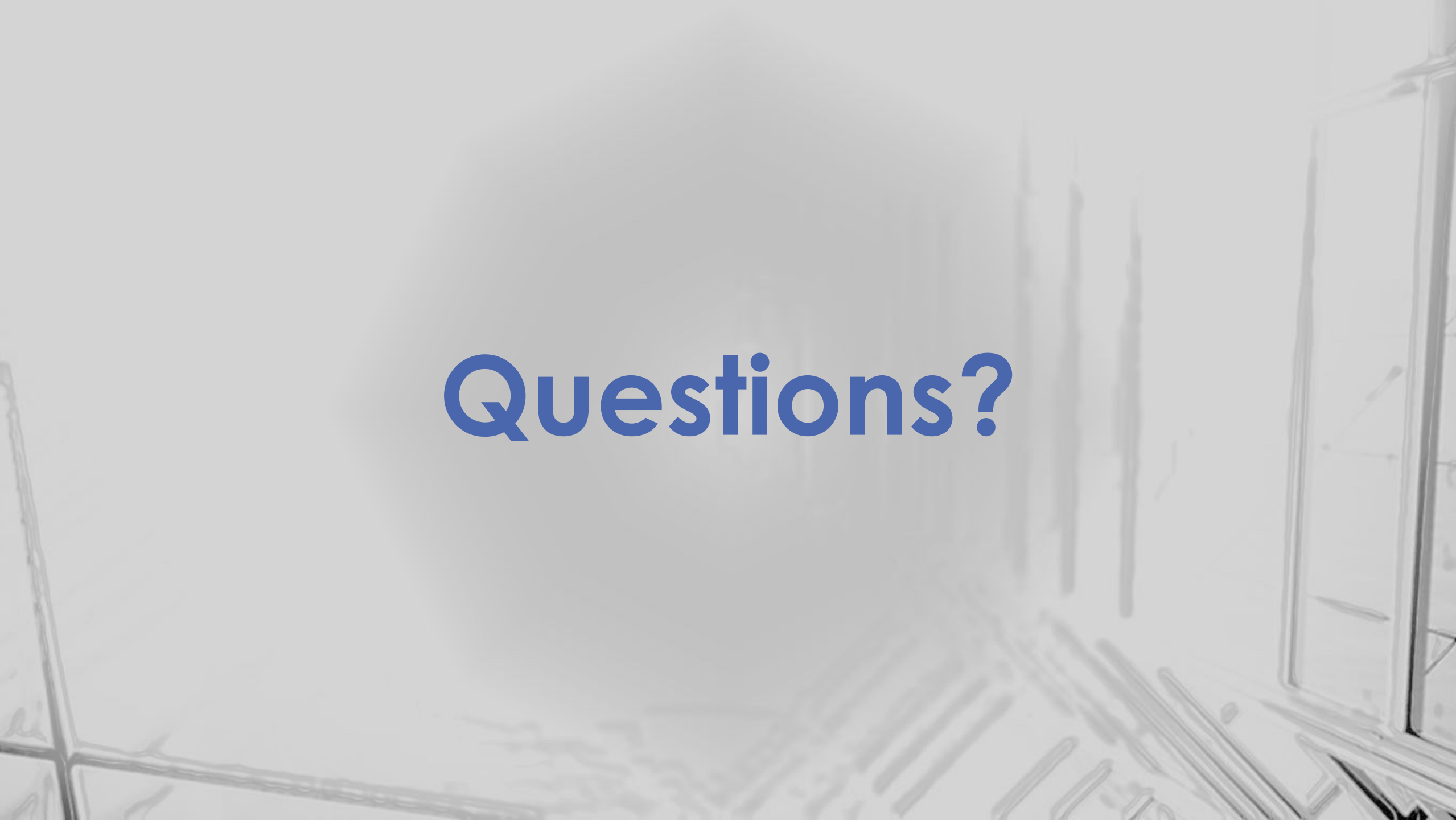
Shifting from
Vulnerability-Based
to
Exploitability-Based
Security



6. Key Actions

Key Actions

1. **Confirm** your Organization's AI Cybersecurity Strategy includes Ethical Considerations and Trustworthy AI strategies
2. Continually **Invest** time to stay current on AI Cybersecurity enhancements which are progressing at an exponential pace
3. Stay **vigilant** in protecting your organization's most valuable assets from AI enhanced cyber threats

The background is a light gray, faded architectural sketch. It shows a perspective view of a building's interior or exterior, with lines suggesting walls, windows, and structural elements. The sketch is done in a loose, hand-drawn style.

Questions?

That's a Wrap

Feel Free to Connect with me on LinkedIn:
<https://www.linkedin.com/in/tina-lampe>
#Global Council for Responsible AI

Thanks for your time and insights today!