



Using AI Effectively with Cybersecurity

Arming and Unleashing your inner Skeptic!



Part I: Grappling with AI Transformation

(Skeptics Welcome)



Quick Poll (show of hands):

How far are we from self aware AI?

- A. Greater than 20 years away
- B. Within the next 20 years
- C. Already happened

Trick question! Everyone already knows that:

SARAH CONNOR: "...The system goes online August 4th, 1997. It [Skynet] becomes self aware at 2:14 a.m. Eastern Time... And Skynet fought back."



In other news...

The New York Times

Google Sidelines Engineer Who Claims Its A.I. Is Sentient



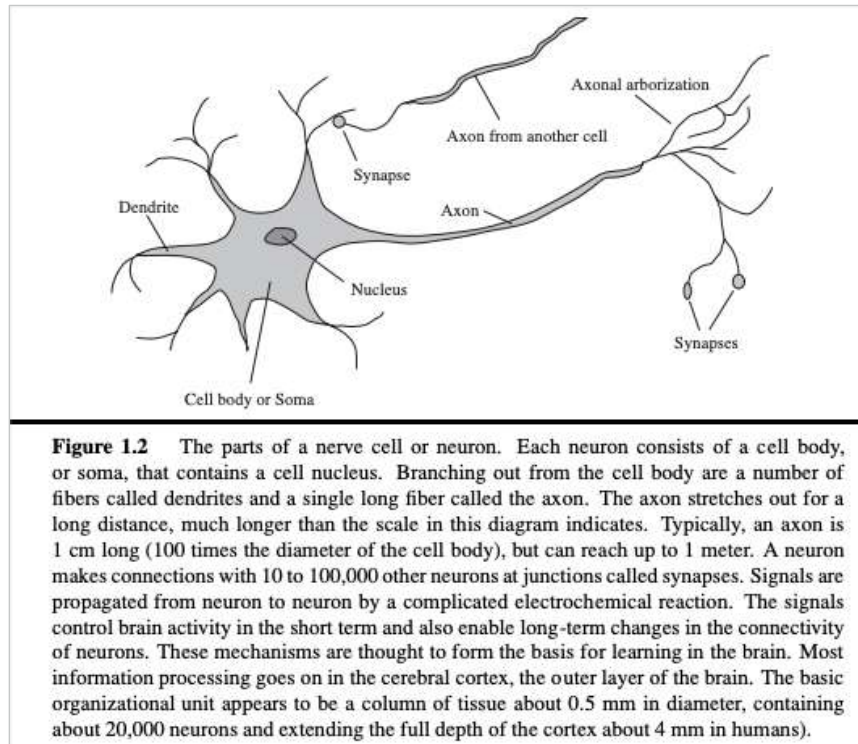
Some artificial intelligence researchers have made optimistic claims about technologies soon reaching sentience, but many others quickly dismiss those assertions. Laura Morton for The New York Times

By Nico Grant and Cade Metz

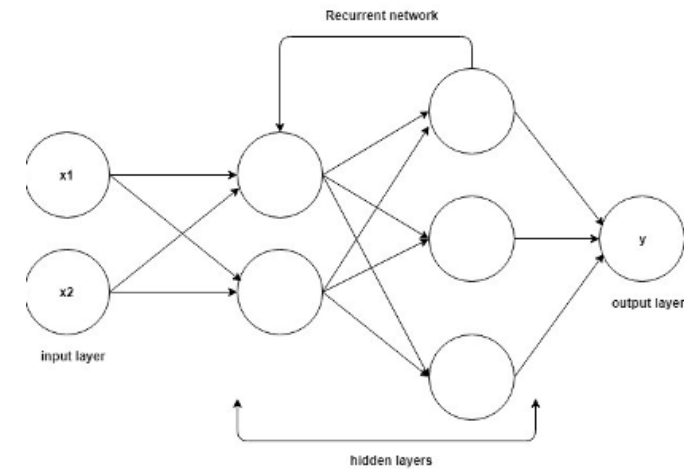
June 12, 2022

<https://www.nytimes.com/2022/06/12/technology/google-chatbot-ai-blake-lemoine.html>

Defining AI: Cognition?



VS.



Artificial Intelligence: A Modern Approach

3rd edition, Russell & Norvig

Defining AI: Not Cognition!

2.2.1 Rationality

What is rational at any given time depends on four things:

- ▼ The performance measure that defines the criterion of success.
- ▼ The agent's prior knowledge of the environment.
- ▼ The actions that the agent can perform.
- ▼ The agent's percept sequence to date.

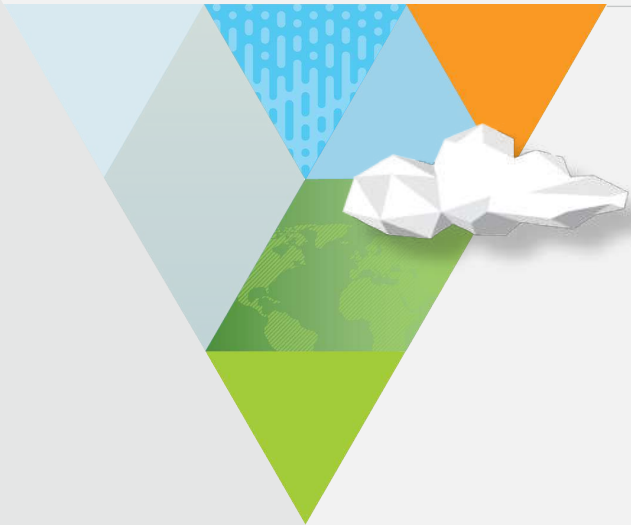
This leads to a **definition of a rational agent**:

For each possible percept sequence, a rational agent should select an action that is expected to maximize its performance measure, given the evidence provided by the percept sequence and whatever built-in knowledge the agent has.

Artificial Intelligence: A Modern Approach

3rd edition, Russell & Norvig

“Modern AI is about machines producing rational outcomes...”



“... not some notion of human-like cognition.”

Defining AI: Differentiating AI Techniques

DALL·E: Creating Images from Text

TEXT PROMPT

AI-GENERATED
IMAGES



Defining AI: Differentiating AI Techniques

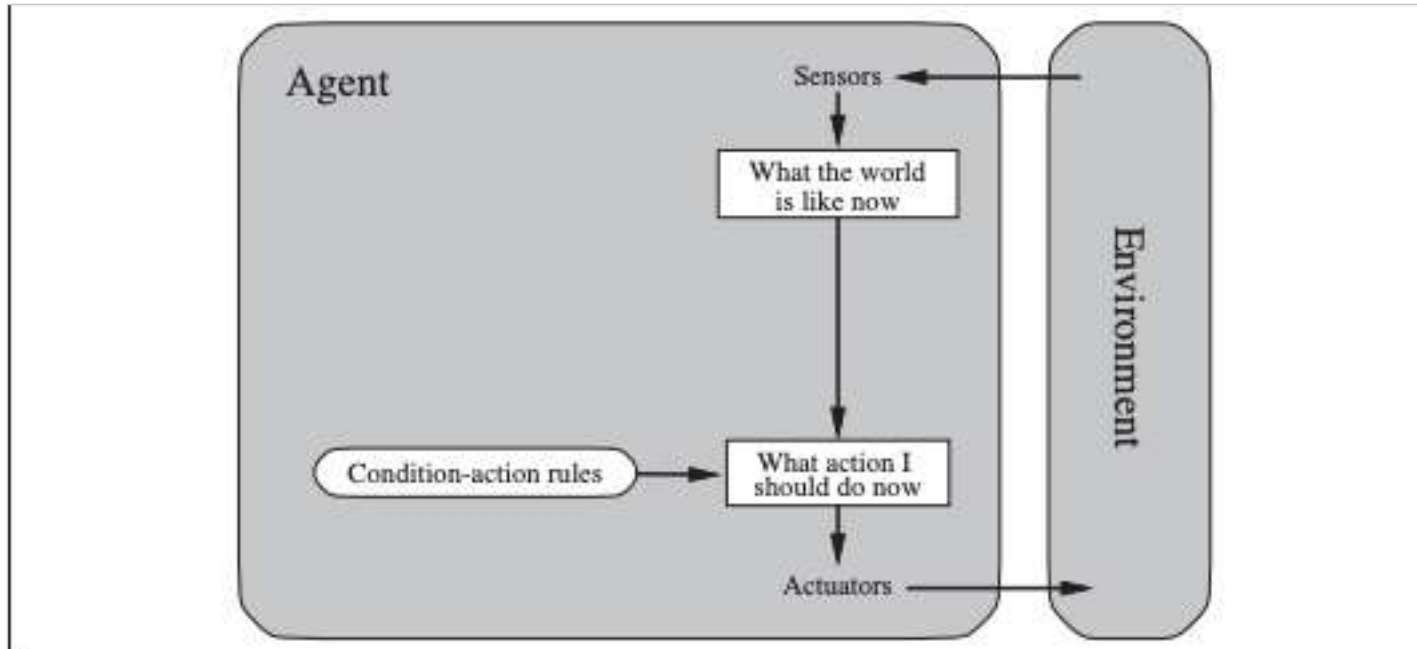


Figure 2.9 Schematic diagram of a simple reflex agent.

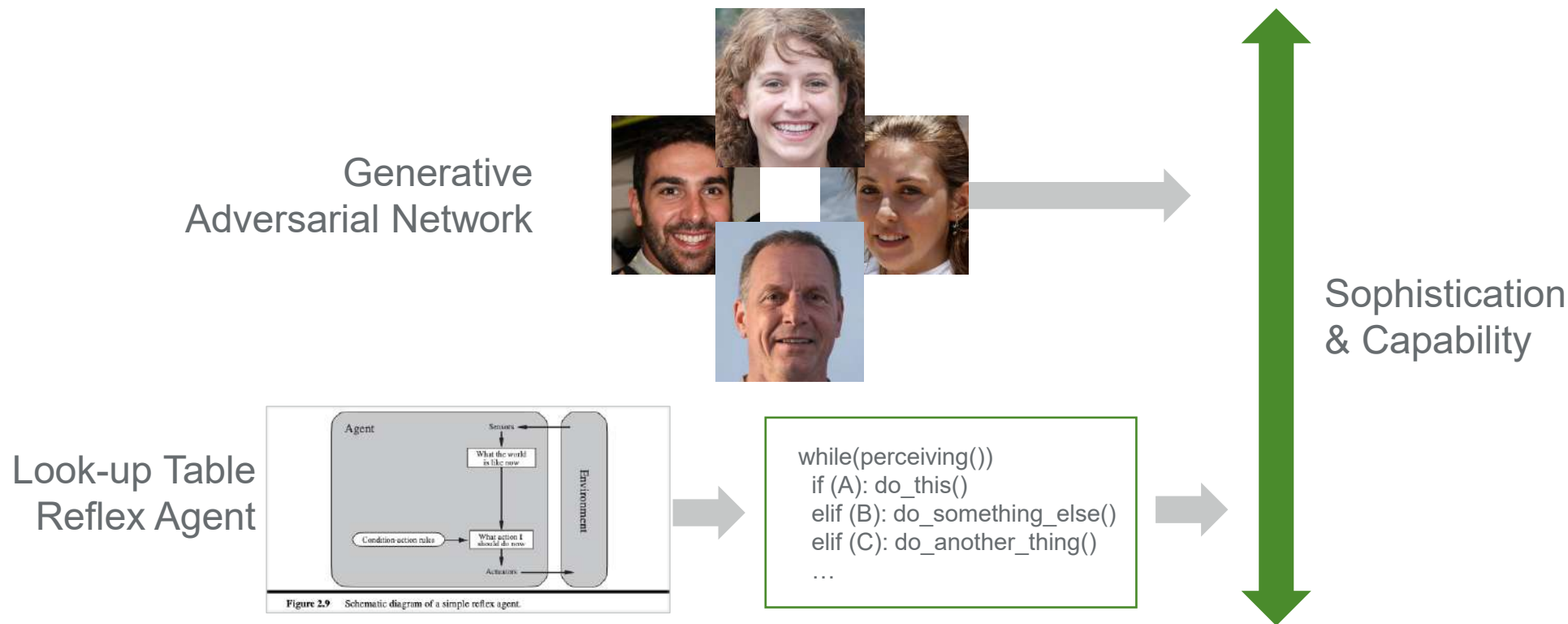
```
function SIMPLE-REFLEX-AGENT(percept) returns an action  
persistent: rules, a set of condition–action rules  
  
state ← INTERPRET-INPUT(percept)  
rule ← RULE-MATCH(state, rules)  
action ← rule.ACTION  
return action
```

Figure 2.10 A simple reflex agent. It acts according to a rule whose condition matches the current state, as defined by the percept.

Artificial Intelligence: A Modern Approach

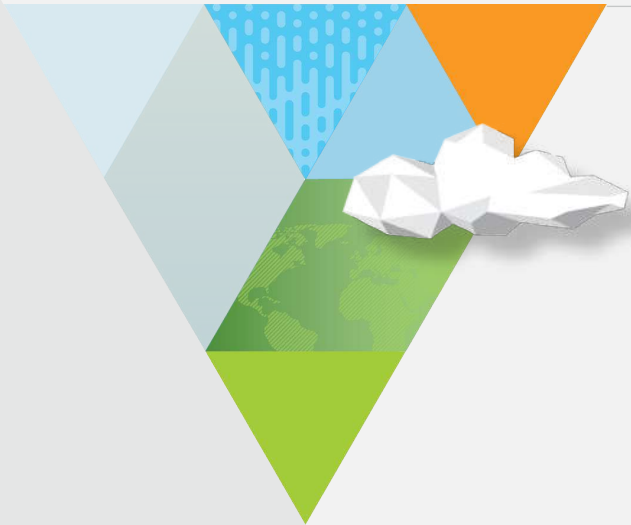
3rd edition, Russell & Norvig

Defining AI: Differentiating AI Techniques



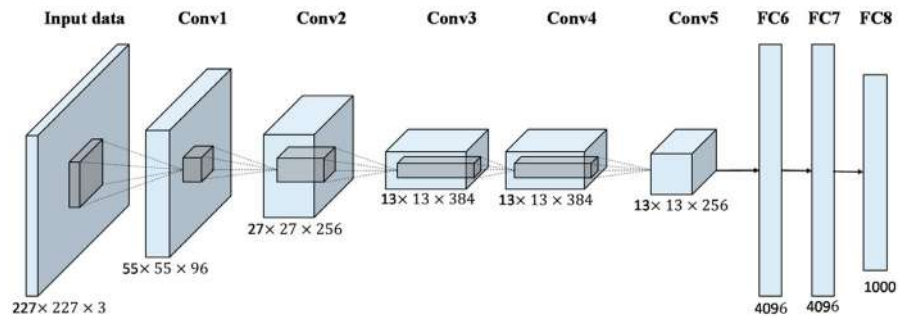
Technique	Is AI?	Useful for Advanced Applications?
Generative Adversarial Network (GAN)	✓	Certainly!
Look-up Table Reflex Agent	✓	Enh... Probably not!

“Useful AI focuses on transformative *outcomes*...”



“...not merely the *presence* of AI.”

Defining AI: AI Threat Surface



= Bus ✓



= Bird ✓

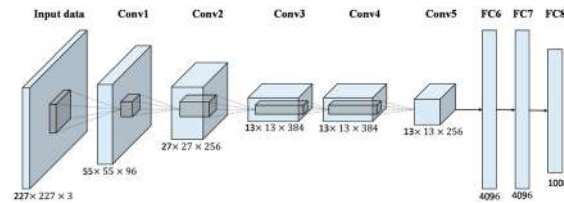


= Temple ✓

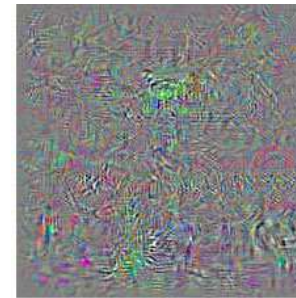
“Intriguing properties of neural networks” (2014)

<https://arxiv.org/abs/1412.6572>

Defining AI: AI Threat Surface



+



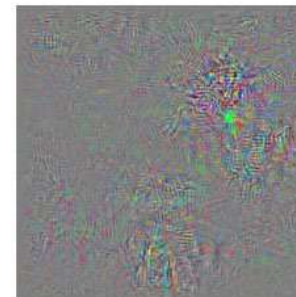
=



= Ostrich



+



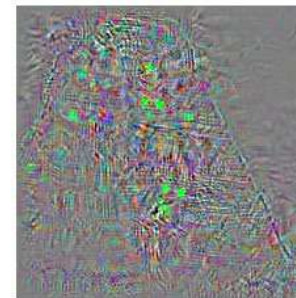
=



= ...Ostrich??



+



=

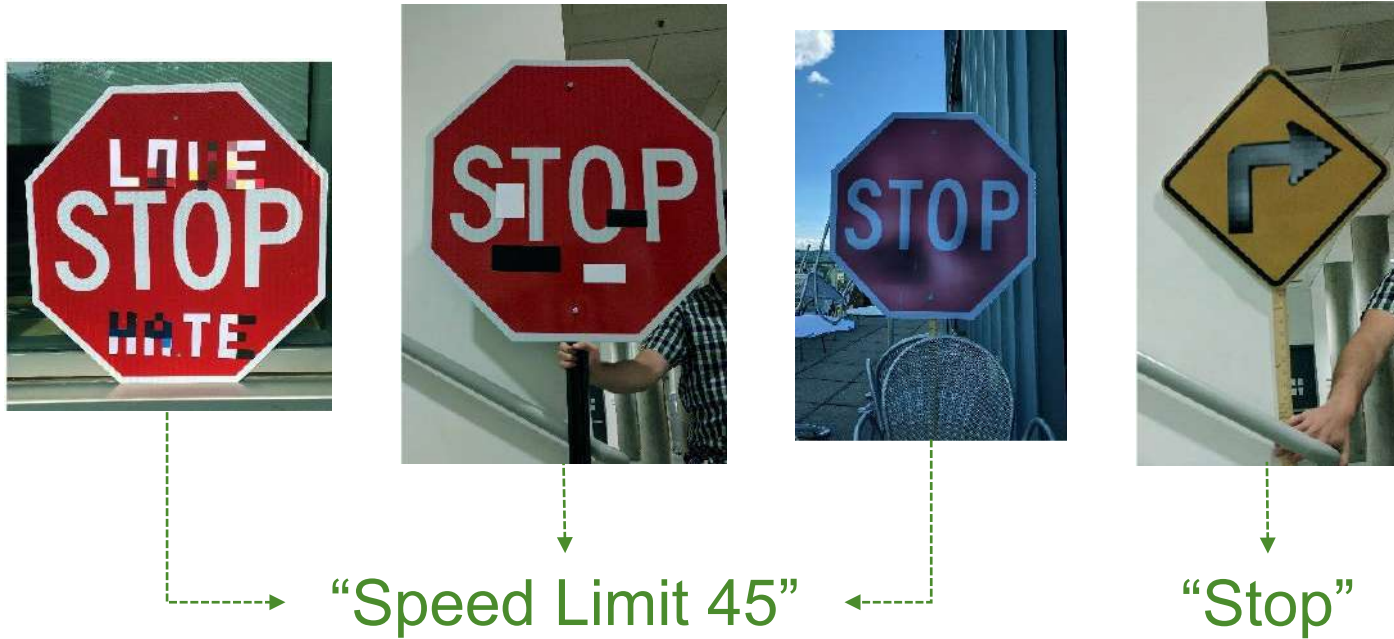


= Ostrich!!?!??

“Intriguing properties of neural networks” (2014)

<https://arxiv.org/abs/1412.6572>

Defining AI: AI Threat Surface



“Robust Physical-World Attacks on Deep Learning Models”-CVPR 2018

Kevin Eykholt^{*1}, Ivan Evtimov^{*2}, Earlene Fernandes², Bo Li³,
Amir Rahmati⁴, Chaowei Xiao¹, Atul Prakash¹, Tadayoshi Kohno², and Dawn Song³

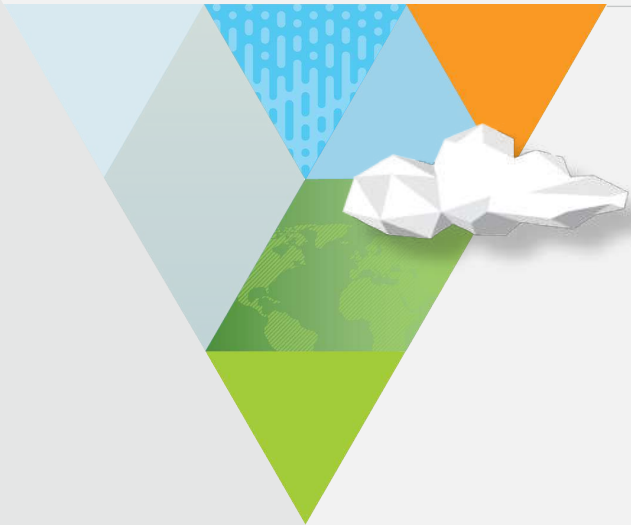
¹ University of Michigan, Ann Arbor

² University of Washington

³ University of California, Berkeley

⁴ Samsung Research America and Stony Brook University

“Transformative AI focuses on success in real-world (unfriendly!) operating environments...”



“...not just controlled, positive use-cases.”

Part II: The Symptoms of Transformative AI



How to think about opportunities for Transformative AI

- ▼ **Start with a Problem Statement** – “*What are we objectively attempting to accomplish?*”
- ▼ **Proceed to a Question** – “*How could an intelligent machine accomplish this more {effectively | efficiently | etc } than a human?*”
- ▼ **Tailor a solution** – “*Embracing No Free Lunch, select a specific AI approach based on the nature of domain and desired outcomes.*”



Performance vs. Generality Trade-off

How to think about opportunities for Transformative AI



*Using a general technique to do **all things** means that single technique will do **all things poorly**.*



*Using a tailored approach to do **one thing** allows that single thing to be done **very well**!*



Performance vs. Generality Trade-off

AI for Outcomes: Practical Challenges

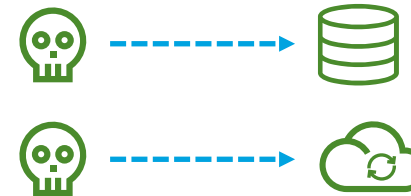
▼ Two practical problems faced by network defenders:

- Detecting Adversary Traffic over encrypted Tunnels:
- Detecting abuse of legitimate-but-compromised credentials:

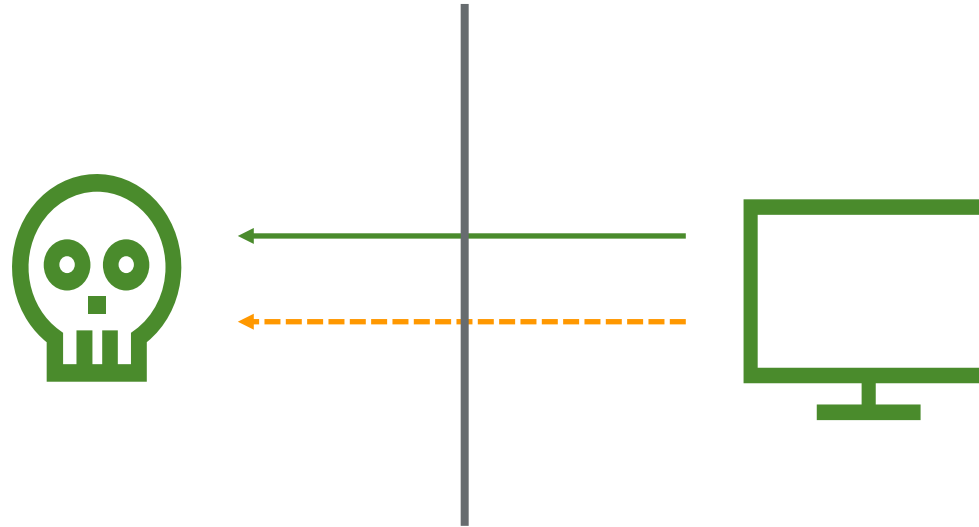
Encrypted Tunnel Detection



Stolen Privileged Network + Cloud Credentials

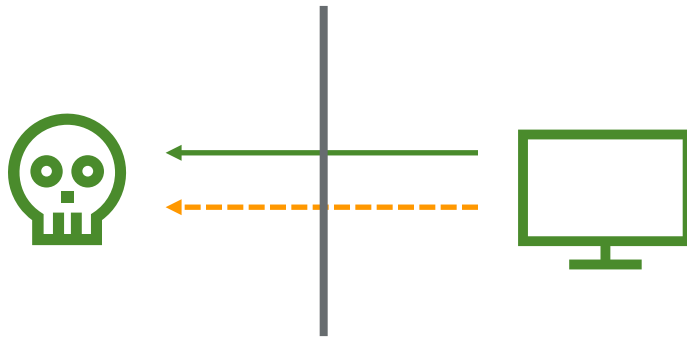


Challenge 1: Detect an HTTPS Tunnel



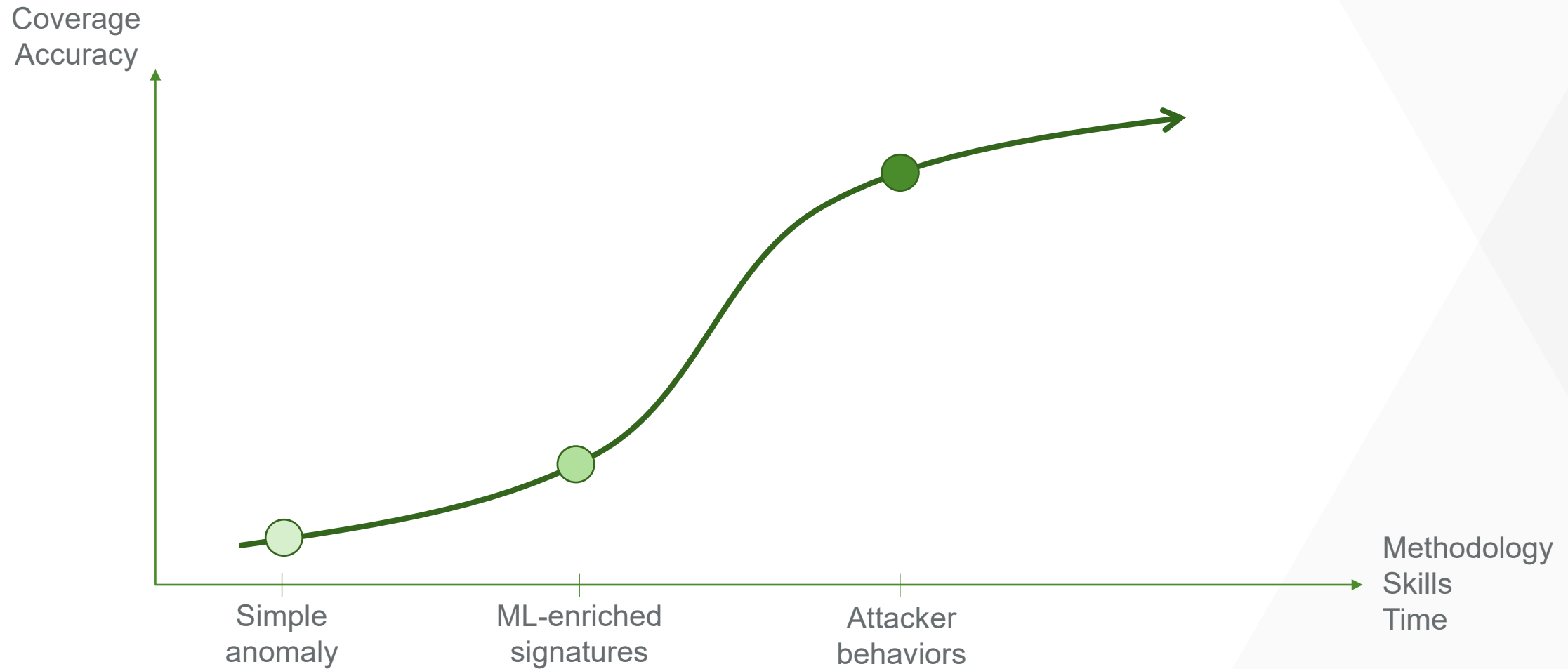
Challenge 1: Detect an HTTPS Tunnel

Core to every APT attack is their C2



- ▼ Designed to evade detection
- ▼ Attackers constantly evolve
- ▼ Benign networks constantly change

Perspective on approaches



Challenge: Detect an HTTPS Tunnel

Durability
Accuracy
Sophistication

Simple anomalies +
static IoCs

Example

Notes

Unusual HTTPS
connection count

FP: user browses more
FN: tunnel conns low relative to user

Challenge: Detect an HTTPS Tunnel

Durability
Accuracy
Sophistication



ML-enhanced signatures

Example

Notes

SSL beaconing to rare destination

Approximation of tunnel -> FP and FN problems

Simple anomalies + static IoCs

Unusual HTTPS connection count

FP: user browses more
FN: tunnel conns low relative to user

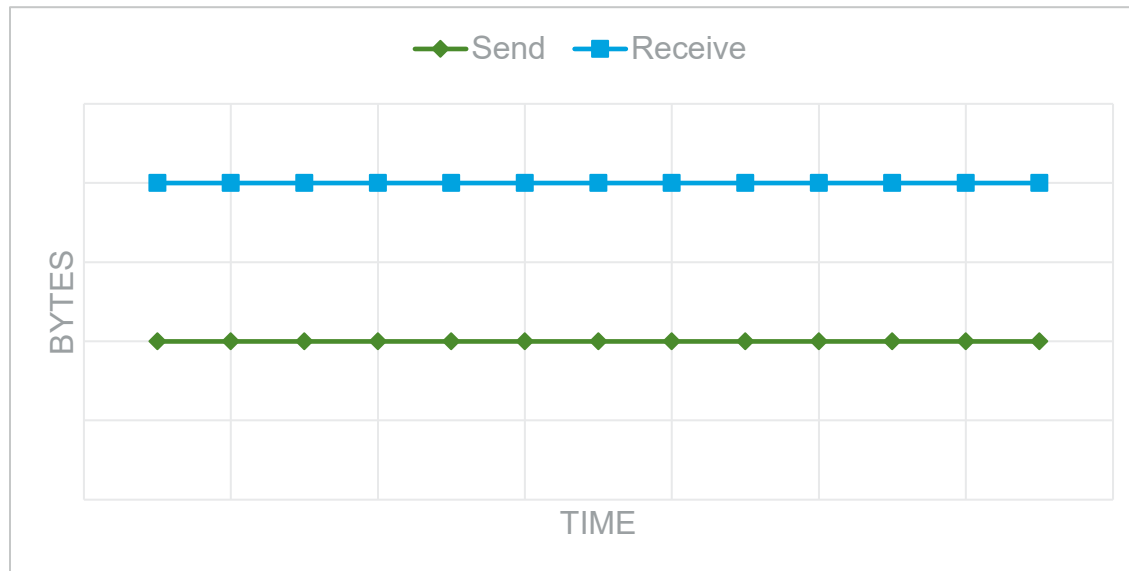
Challenge: Detect an HTTPS Tunnel

Durability
Accuracy
Sophistication

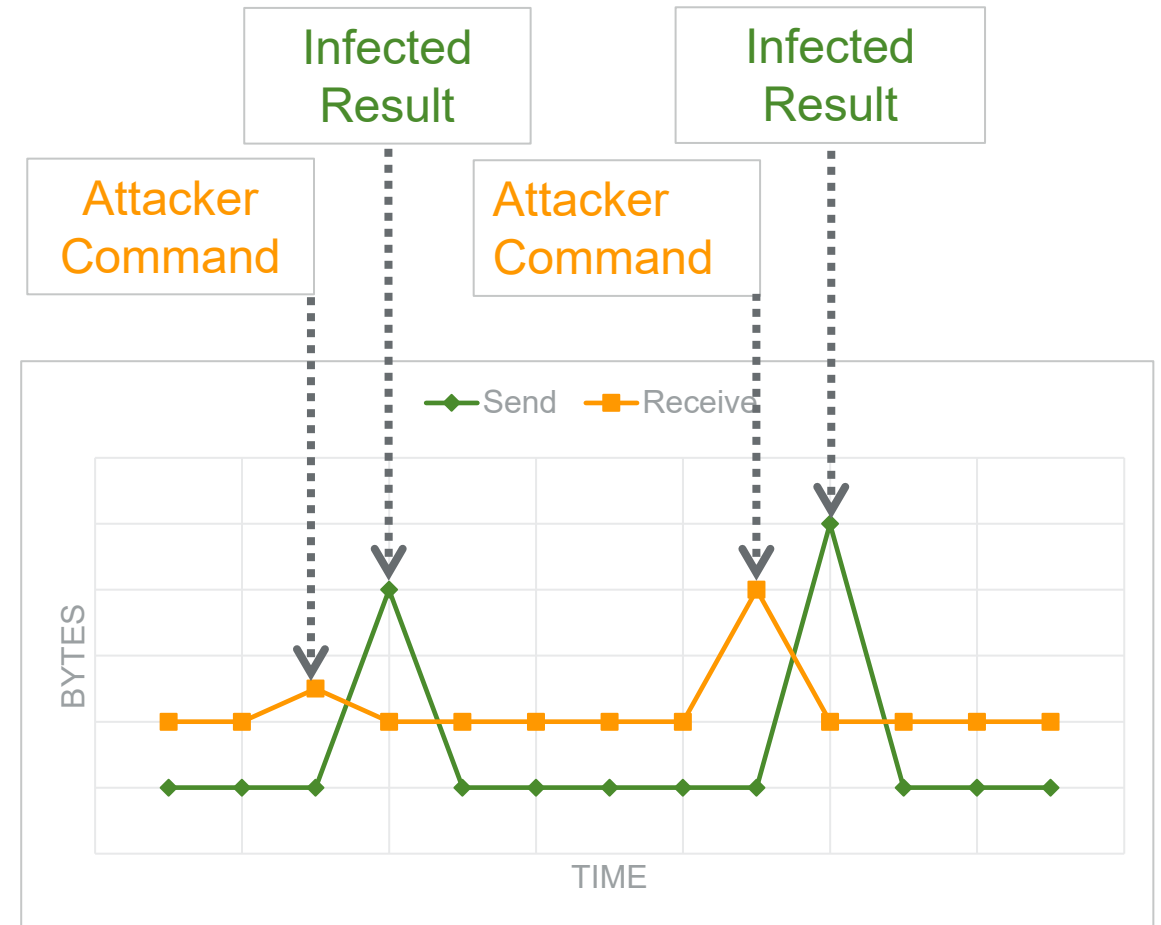
	Example	Notes
AI model	HTTPS Tunnel detector	Directly detect behavior of interest in an accurate way
ML-enhanced signatures	SSL beaconing to rare destination	Approximation of tunnel -> FP and FN problems
Simple anomalies + static IoCs	Unusual HTTPS connection count	FP: user browses more FN: tunnel conns low relative to user

Visible control in the data

Beacons and long running connections



Benign Traffic



Attacker Traffic

Hidden HTTPS Tunnel Model

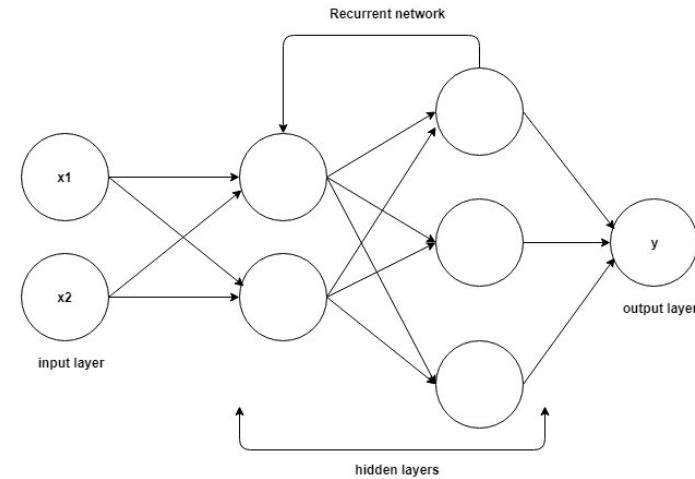
1000s of labeled
tunnel samples

- ▼ Many tools
- ▼ Many uses

Normal HTTPS
traffic from across
environments

Time series data.
Sub-second data
transfer patterns.

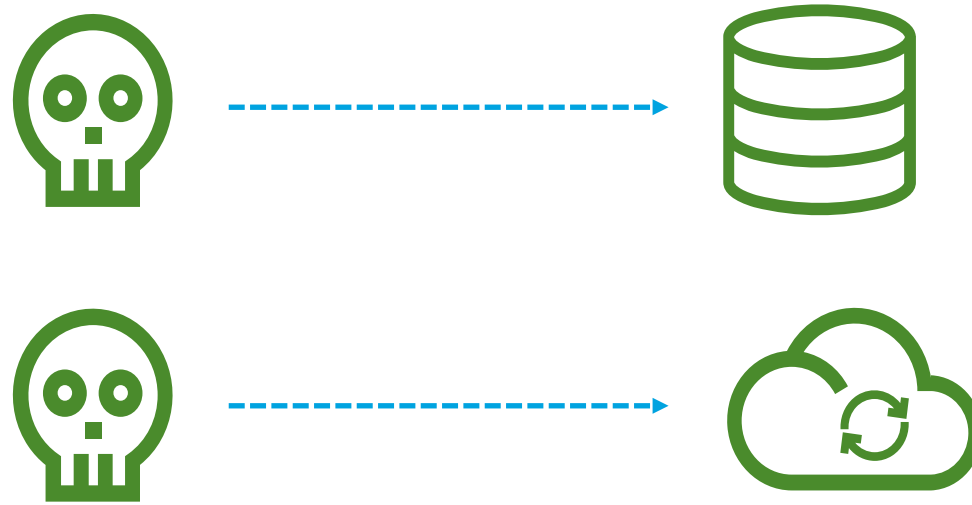
Train



Directly detects
behavior of
tunneling.
No whitelists.
No blind spots.

Deep Learning:
LSTM Recurrent Neural Network

Challenge 2: Detecting the abuse of privilege credentials



Challenge 2: Detecting the abuse of privilege credentials

Privilege accounts are high priorities for attacker



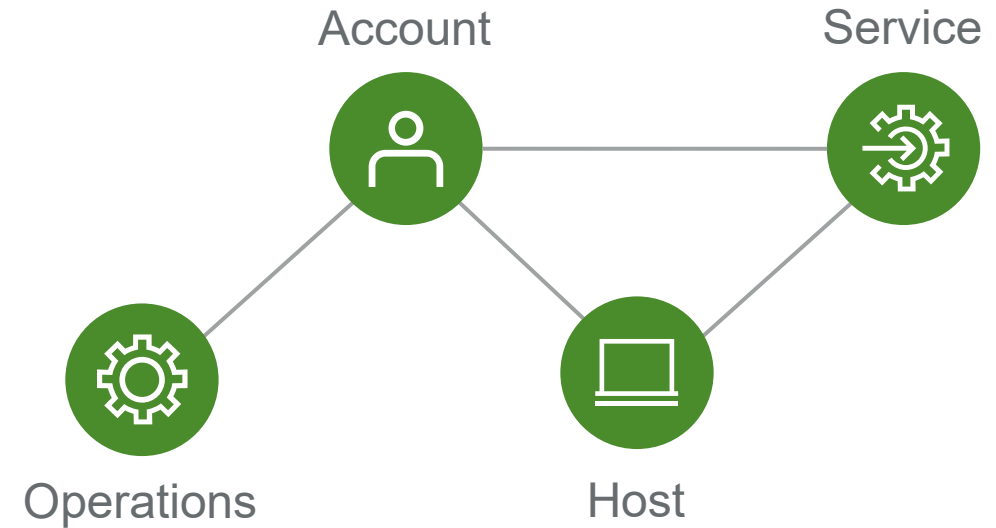
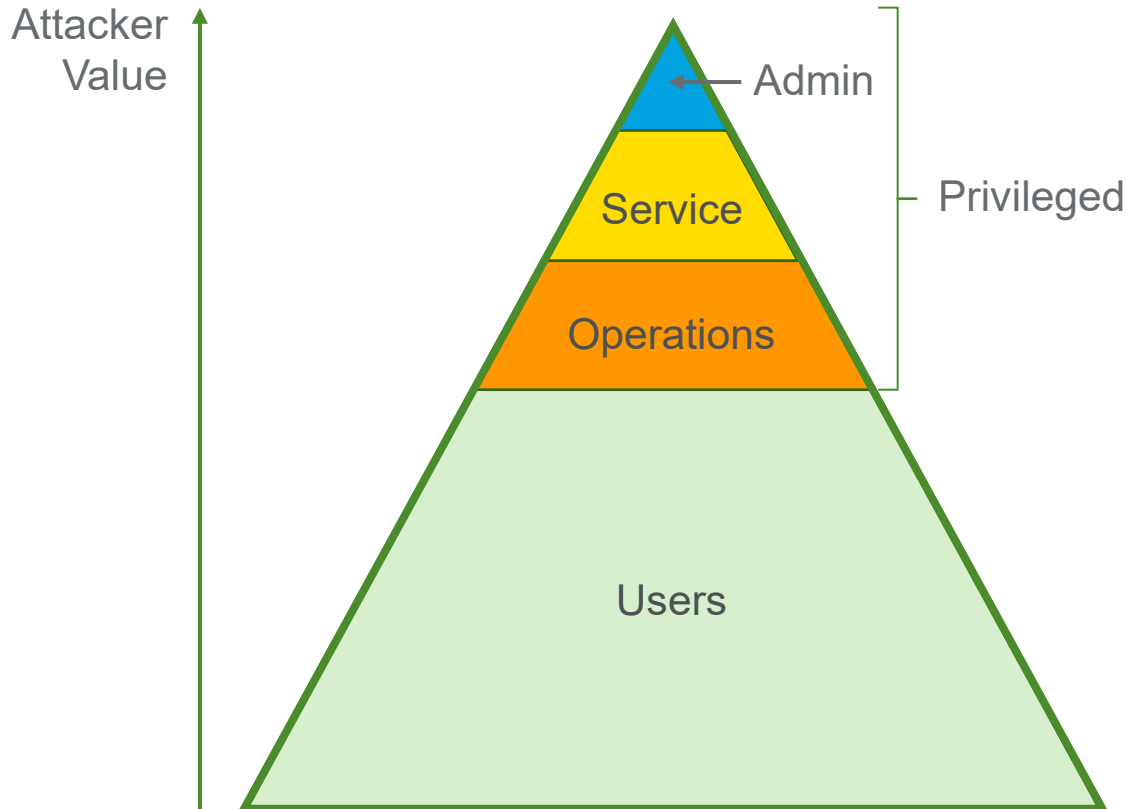
- ▼ Access to both **network** and **cloud**
- ▼ By definition, actions are **allowed** to happen
- ▼ Abnormal **is** normal

Perspective on approaches

Durability
Accuracy
Sophistication

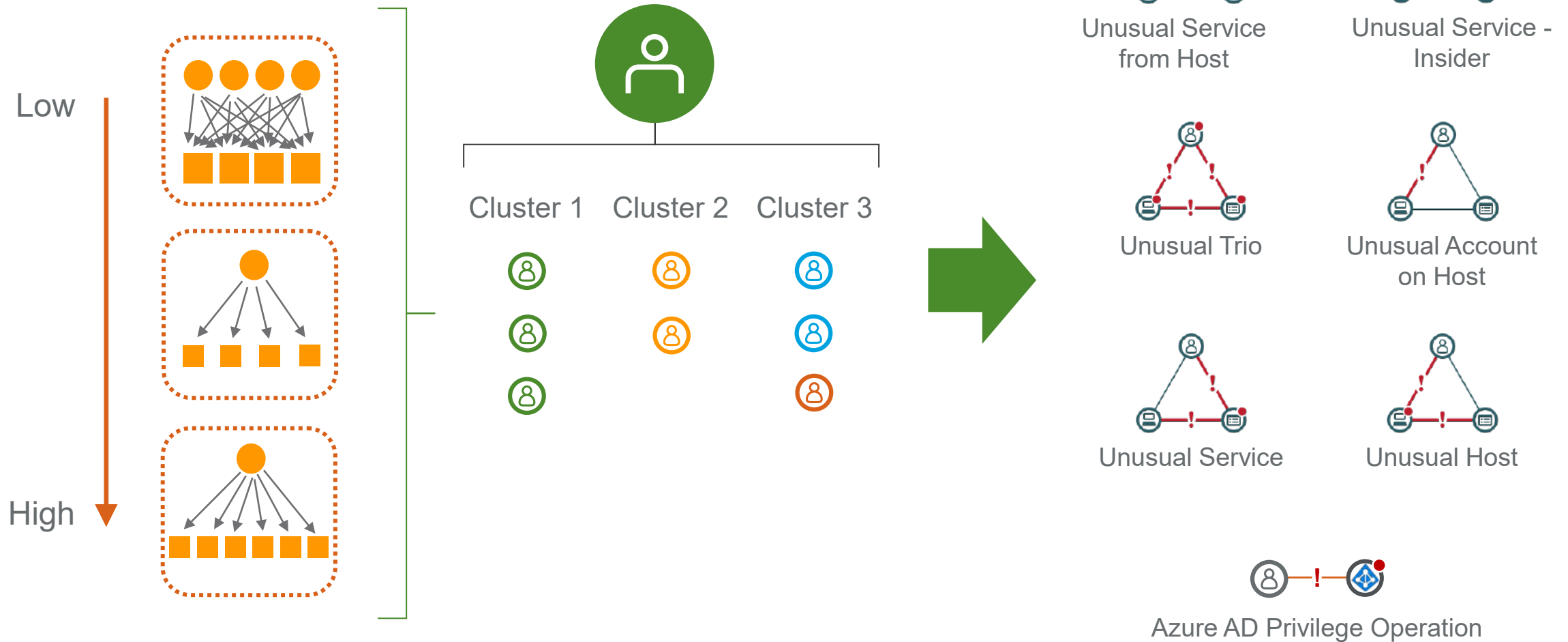
	Example	Notes
AI model	Privilege aware anomaly	Detect how attackers abuse of credentials
“Rare” access	First access if it is used by less than 5 others	FP: New functionality FN: Attacker access to common service
First access	First time an account accesses something	FP: New functionality in org FN: Access on different host

Not all access is valuable



- ▼ Map relationships
- ▼ Observe and learn **true** privilege
- ▼ Detect useful anomalies

Privilege Anomaly Models



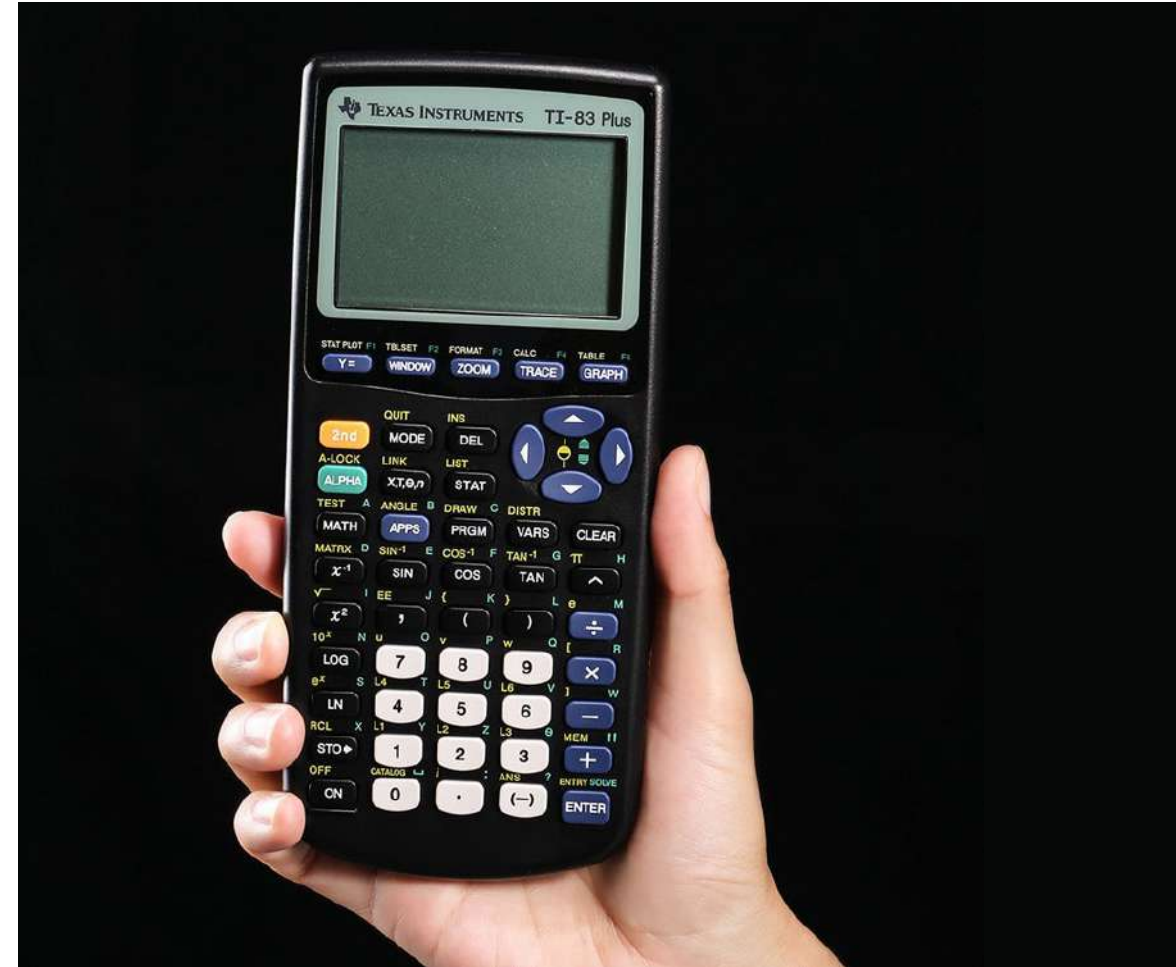
Part III: AI Transformation for Security

Humans and Machines



But first, let's talk about Calculators...

Did calculators replace
mathematicians?



AI Teams: Human / Machine Independent Task Excellence

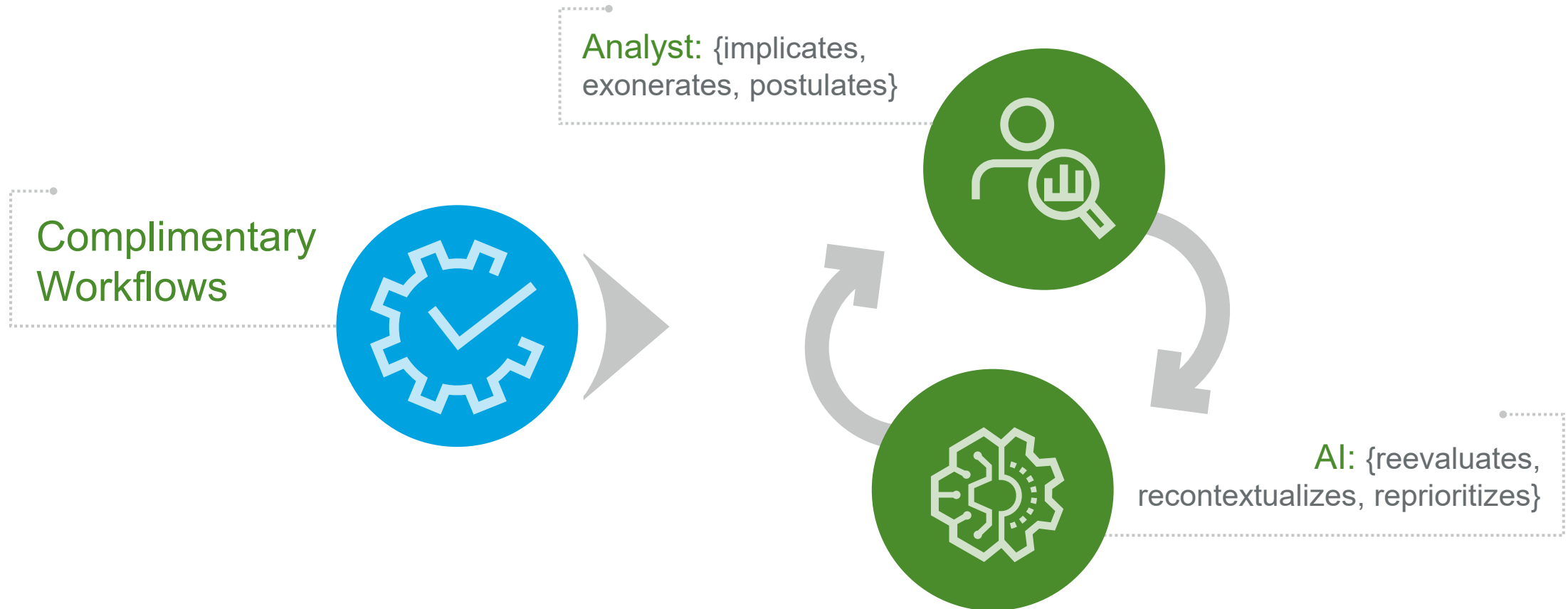


- ▼ Ethics / Culture
- ▼ Operational Ambiguity
- ▼ Abstract Planning / Reasoning
- ▼ Judgment

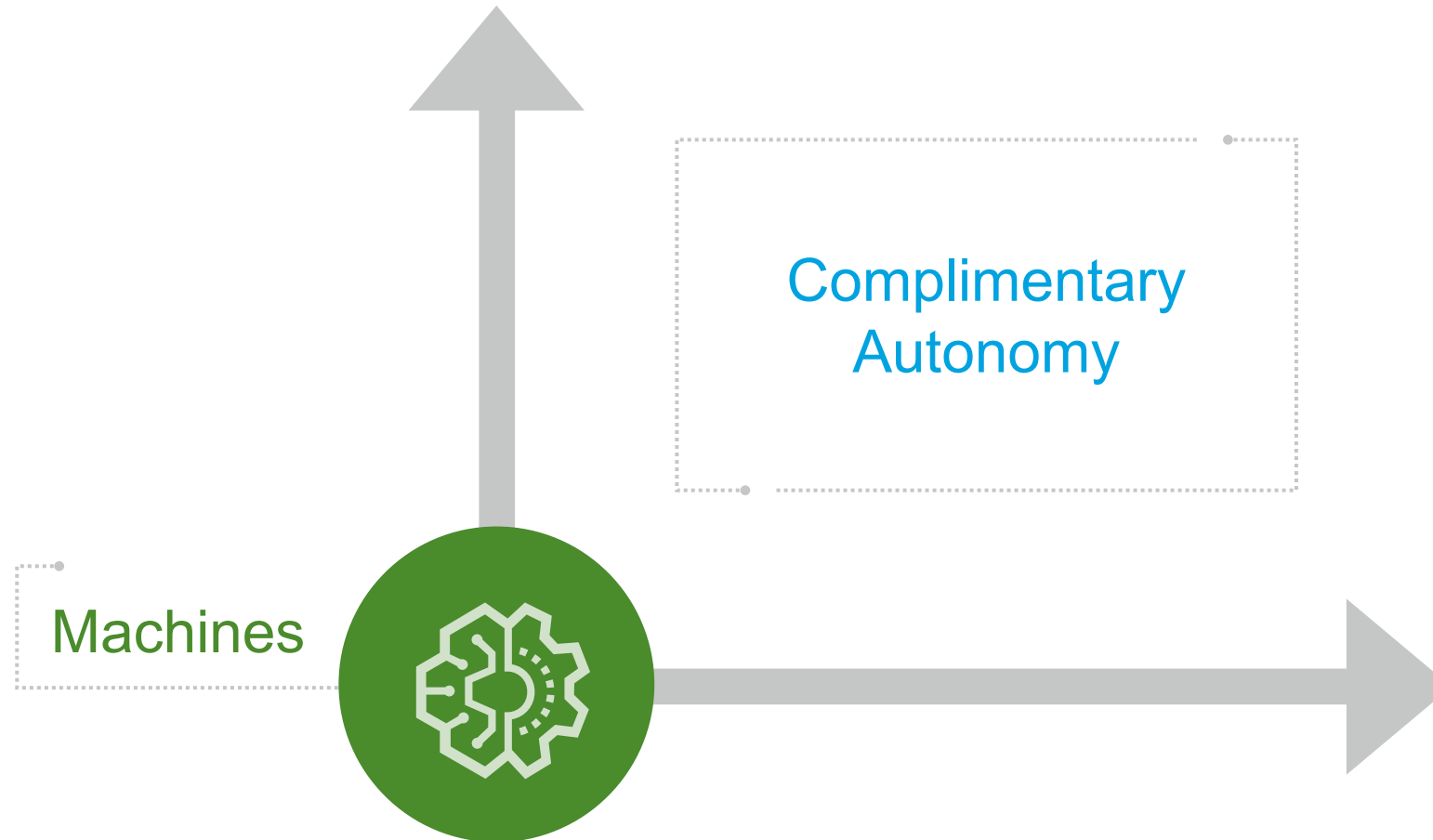
- ▼ Machine Speed
- ▼ Machine Scale
- ▼ Machine Complexity



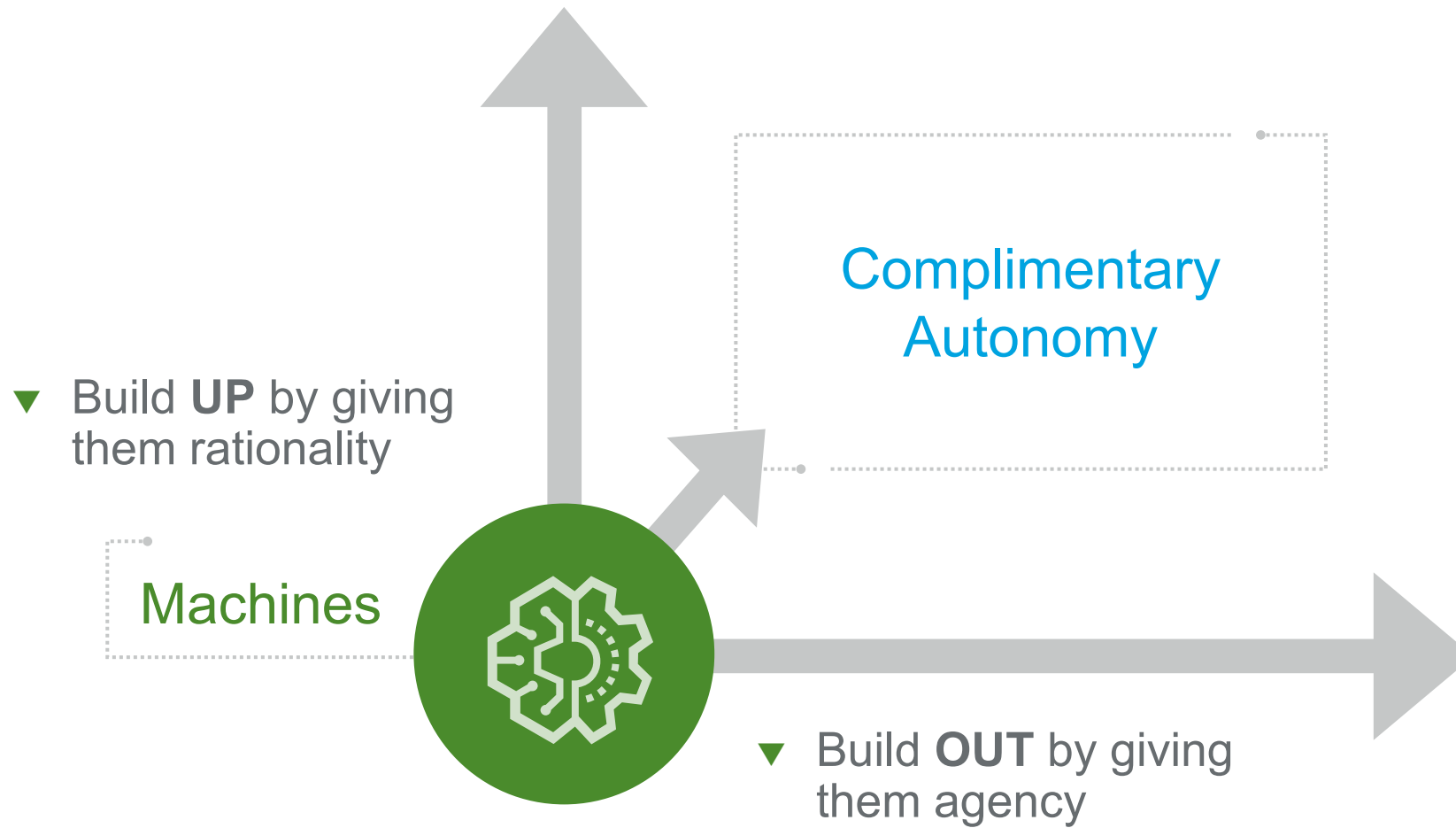
AI Teams: Human / Machine Task Excellence



Effective AI: Vertical / Horizontal Capabilities



Effective AI: Vertical / Horizontal Capabilities



Concluding Thoughts



Take-aways 1:

- ▼ AI Skepticism is *completely rational* in the face of today's hype cycle.
- ▼ Evaluating AI should focus on:
 - Outcomes
 - Improvements vs. existing state of the art
 - Maintaining resilience across operational environments

Take-aways 2:

▼ Effective AI Transformation:

- starts with a *problem*
- proposes an *improvement*
- purpose-builds an *outcome for that problem*

Take-aways 3:

▼ AI Security Transformation:

- Enables *humans* to be better *humans*
- Enables *machines* to be better *teammates* through:
 - Rational outcomes
 - Environmental Agency

THANK YOU!!

Q&A

