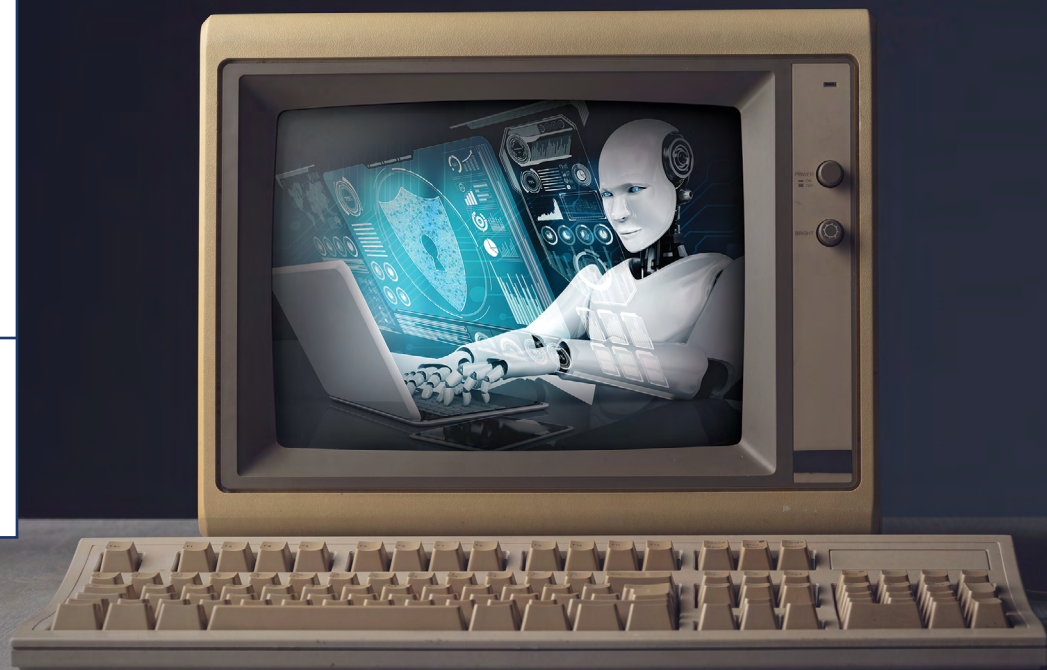




Ann Caracristi Institute
For Intelligence Research

RESEARCH FELLOWSHIP



RESEARCH MONOGRAPH

Perceptions of Artificial Intelligence/Machine Learning in the Intelligence Community A Systematic Review of the Literature

Adrian Wolfberg, Ph.D.

Research Fellow, National Intelligence University, 2020-21

The views expressed in this research monograph are those of the author and do not reflect the official policy or position of the National Intelligence University, Office of the Director of National Intelligence, or any other U.S. Government agency.

Perceptions of Artificial Intelligence/ Machine Learning in the Intelligence Community:

A Systematic Review of the Literature

Adrian Wolfberg, Ph.D.

NATIONAL INTELLIGENCE UNIVERSITY
PUBLISHED FALL 2022

The views expressed in this research monograph are those of the author and do not reflect the official policy or position of the National Intelligence University, the Office of the Director of National Intelligence, or any other U.S. Government agency.

Abstract

Many voices have expressed their concerns about the state of artificial intelligence (machine learning) in the U.S. Intelligence Community. This study uses a systematic review methodology to collect, analyze, and synthesize the reasons for these concerns, the problems raised, and the solutions offered. A systematic review aims to provide an evidence-based management approach to identify what is known about a topic and provide practitioners with information to decide what action to take in the future. Using a systematic review methodology, this study appears to be the first of its kind and captures a holistic understanding of issues related to artificial intelligence in the Intelligence Community. The theory of organizational attention is used as a theoretical lens because the findings are very complex regarding the breadth and depth of the problems raised and the solutions offered. They directly affect what and how Intelligence Community decisionmakers focus their attention on and react to the findings. The concerns raised in the literature address issues related to decisional and behavioral factors leading to artificial intelligence usage. The problems identified are categorized into those internal to the Intelligence Community, those external to the Intelligence Community, and those unique to artificial intelligence. The solutions address what knowledge gaps need to be filled and what implementation needs should be satisfied. A conceptual model is proposed to help absorb and operationalize the findings from this study. While artificial technology is a technology, almost everything in this study is about psychology, social psychology, and organizational behavior, i.e., it is a story about people.

Table of Contents

ABSTRACT	3
INTRODUCTION	9
Issue	9
Research Question	9
Scope	9
Purpose	10
Relevance to the Intelligence Community	10
RESEARCH METHODOLOGY	11
Conceptual Framework	11
Key Questions	12
Research Design	13
<i>Systematic Review Method</i>	13
Data Collection Strategy	14
<i>Search Details</i>	16
Analytic Strategy	17
FINDINGS	19
Motivation	19
<i>Decisional Issues</i>	19
AI Considerations	19
Non-AI Considerations	20
<i>Behavioral Issues</i>	21
AI External Concerns	21
AI Internal Concerns	22

Problems	23
<i>Internal to IC Problems</i>	23
IC Culture Constraints	23
Outdated Theory of Intelligence	25
Perceptions	26
System-Level Constraints	28
<i>External to IC Problems</i>	29
External Threats	29
Policy Constraints	31
Digital Environment	31
<i>Nature of AI/ML</i>	32
Knowledge Constraints	32
Limitations of AI/ML	34
Solutions	36
<i>Fill Knowledge Gaps</i>	36
Decide What AI/ML Should Do For... ..	37
Scholarly Research	39
Leadership Visioning	40
<i>Address Implementation Needs</i>	42
Strategy Development	42
Team Development	44
Community Focus	45
Summary of Findings	46
<i>Problems and Solutions</i>	47
CONCLUSION AND IMPLICATIONS	49
Multilevel Nature of Problems	49
Solutions at Each Level	50
Organizational Attention	51
Ways to Lead	51
Summary of Conceptual Framework	52
Recommendations	52
Limitations	53
Concluding Comments	53

ENDNOTES	55
----------------	----

REFERENCES	69
------------------	----

LIST OF FIGURES

Figure 1. Theory of Organizational Attention applied to the IC	12
Figure 2. Three Subordinate Review Questions about AI/ML within the Intelligence Community ...	13
Figure 3. Search Flow of AI/ML Systematic Review	15
Figure 4. Distribution of Final 41 Articles by Publication Year	16
Figure 5. Types of Publication Venues	16
Figure 6. Author's Job at Time of Publication	17
Figure 7. Motivation for Why AI/ML Issues are a Concern for the IC	19
Figure 8. Motivational Decisional Issues	19
Figure 9. Motivational Behavioral Issues	21
Figure 10. Problems of AI/ML for the IC	23
Figure 11. IC Culture Constraints Problems	24
Figure 12. Outdated Theory of Intelligence Problems	25
Figure 13. Perceptions of AI/ML Problems	26
Figure 14. System-Level Constraints Problems	28
Figure 15. External Threats Problems	30
Figure 16. Policy Constraints Problems	31
Figure 17. Digital Environment Problems	32
Figure 18. Problems Associated with Knowledge Constraints	33
Figure 19. Problems Associated with Limitations of AI/ML	35
Figure 20. Solutions for AI/ML for the IC	37
Figure 21. Making AI/ML Decisions Solutions	37
Figure 22. Scholarly Research Solutions	39
Figure 23. Leadership Visioning Solutions	41
Figure 24. Strategy Development Solutions	42
Figure 25. Team Development Solutions	44
Figure 26. Community Focus Solutions	45

Figure 27. Answers to Review Questions	46
Figure 28. AI/ML in the IC as a Multilevel Human Problem Set	50
Figure 29. Conceptual Framework for Addressing AI/ML in the IC	52

LIST OF TABLES

Table 1. Review Questions and Coding Results	18
--	----

Introduction

Issue

This study is part of a broader project to create a vision for the future of the national security intelligence mission.¹ This study is a precursor to such a vision as it suggests what is known about the overall question and the current state of knowledge of problems and solutions within the Intelligence Community (IC) today.

Research Question

The research question is, what is the state of knowledge about problems and solutions in today's Intelligence Community with AI/ML?

Scope

This study focuses on artificial intelligence and machine learning in the Intelligence Community. The terms for these technologies are often used interchangeably in practice, hence they are combined as AI/ML.

Artificial intelligence can be defined as technologies that alter themselves, such as machine learning, statistical techniques using algorithms, and deep learning, a more sophisticated form of machine learning.² Artificial intelligence can also be defined as technology that solves problems that a human typically does.³ However, there is no standard definition for artificial intelligence.⁴ The National Security Commission on Artificial Intelligence recently defined artificial intelligence as "...technologies that solve tasks requiring human-like perception, cognition, planning, learning, communication, or physical action; and technologies that may learn and act autonomously, whether in the form of software agents or embodied robots."⁵

Machine learning benefits users based on how the technology has been trained. Supervised training means a priori data has been categorized by humans relevant to the purpose of the technology. Unsupervised training means data does not have to be categorized a priori, and deep learning uses algorithms called neural networks.⁶ For the most part, machine learning solves problems where a priori knowledge exists about the data complexities and how various data interact.⁷ Many articles use the term artificial intelligence without further specification. In contrast, other articles focus on machine learning, the most used advanced technology in the Intelligence Community.

Purpose

The purpose for answering the research question is to ground decisionmakers in the reality of the current state of knowledge of AI/ML. This knowledge can then lead to discussions and decisions about how the IC can evolve into a more effective, relevant, and sought-after source of information so that.

Relevance to the Intelligence Community

This study is focused on AI/ML as it applies to the IC. The different perspectives and knowledge about this topic can provide a knowledge baseline to build ideas about its future for the IC.

Research Methodology

Conceptual Framework

Organizational behaviorist William Ocasio's attention-based view of the firm theory provides the theoretical lens for this study from an individual and organizational perspective.⁸ Ocasio's process-based theory is based on three premises at the individual level of analysis. First, decisionmakers make decisions based on what they attend to. Second, what decisionmakers concentrate on depends on the context of the situation. Third, the context of the situation decisionmakers find themselves in depends on the organization's rules, resources, and social relationships. What we pay attention to influences how problems and answers are handled.

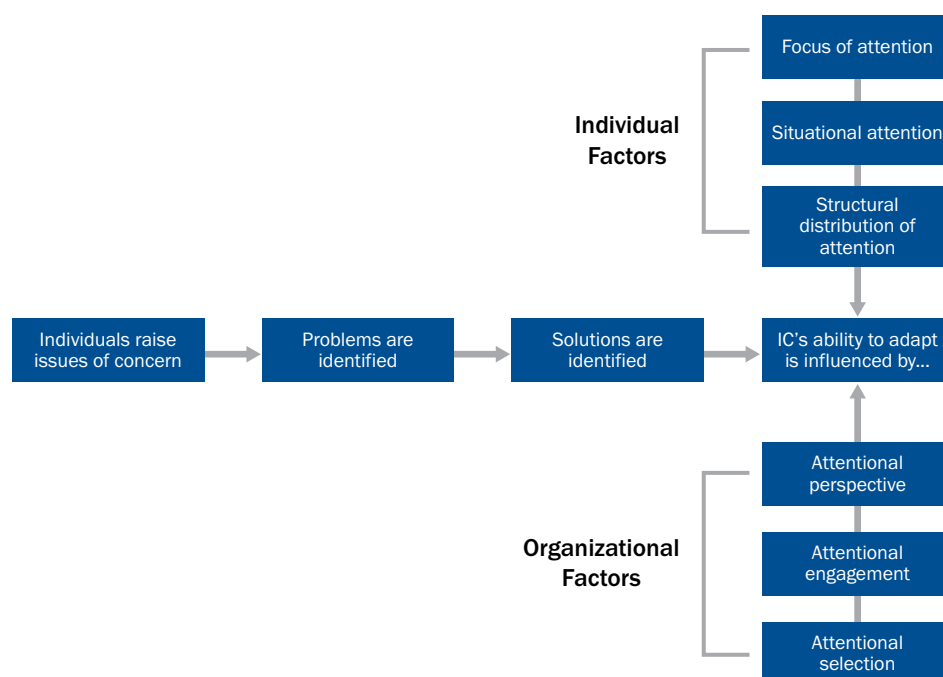
Ocasio asserts that how an organization pays attention to problems shapes the organization's ability to make changes, an adaptation-based theory.⁹ There are three ways that attention shapes the ability of organizations to adapt. First, attentional perspective is how the organization's strategy influences what a decisionmaker pays attention to. Second, attentional engagement is the time, energy, and effort decisionmakers participate in responding to issues, problems, and answers. Third, the attentional selection is the patterns of historical outcomes from paying attention to issues, problems, and solutions.

This study explores the articles discussing AI/ML in which writers implicitly direct their concerns to decisionmakers and organizations. Individuals, whether internal and external to the IC, identify areas of concern they believe the IC needs to address to make it more relevant and valued by its customers. These individuals are motivated to raise issues because they either study the IC, support the IC, or have had experience being a member of the IC. They state the problems and offer solutions in various venues: academic journals, industry concept papers, think tank papers, conference material, etc. Intelligence Community decisionmakers can then consider how the issues, problems, and solutions might be treated. However, if there is a plethora of problems and solutions, then how does the IC absorb them? Which problems or solutions should a decisionmaker pay attention to? Which ones are ignored because they have failed to reach the attention of a decisionmaker?

The theory of organizational attention suggests that the IC's ability to adapt to a new environment is directly affected by the process and engagement types of attention used by IC decisionmakers. The purpose and outcome of this study will not answer the question, "how does the IC absorb information about AI/ML and act upon it?" Additionally, it will provide a reasonably accurate way of organizing the broad spectrum of problems and solutions so that future questions can be asked. One future question may be, how will the IC use this information to adapt to AI/ML?

Figure 1 represents how the theory of organizational attention applies to the IC. Issues are raised to identify problems within a topic, and solutions are offered based on the defined problems. Individuals who have identified problems and solutions are seeking the attention of IC decisionmakers to change the IC for the better. The IC's ability to change partly depends on the decisionmaker's ability to focus their attention, as seen by individual factors in the top half of Figure 1. It is affected by the attention-related behaviors of their organization, in the bottom half of Figure 1, organizational factors. Once the issues, problems, and solutions have been identified, the IC can then consider how this information can be used to make changes to improve itself.

Figure 1. Theory of Organizational Attention applied to the IC



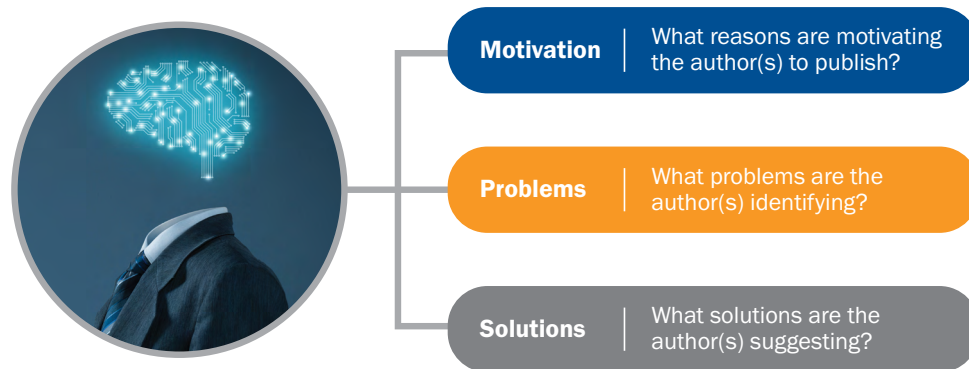
Key Questions

The study used three subordinate review questions to organize the literature analysis, as shown in Figure 2. The term review question is used intentionally to differentiate it from a “research question.” A research question is used when an individual collects data to answer their question (called “primary research”). In contrast, a review question is used when an individual compiles other authors’ finished products, often called “secondary research”. This study captures others’ work and is, therefore, secondary research.

The first subordinate review question (RQ1) asks about someone’s motivation for writing about AI/ML in the IC, “What reasons motivate someone to publish their knowledge and experience?” In other words, what is their focus of attention? The second subordinate review question (RQ2) centers on the stated problems in their publication about AI/ML in the IC, “What problems are identified about AI/ML in the IC?” The author’s focus of attention serves as a lens through which problems are identified. The third subordinate

review question (RQ3) focuses on the solutions offered, “What solutions are suggested that could improve AI/ML in the IC?” Figures 7 through 26 show how each review question is answered in more detail.

Figure 2. Three Subordinate Review Questions about AI/ML within the Intelligence Community



Research Design

This study adopts many, but not all, features from the systematic review methodology, which uses a disciplined approach to find, analyze, and synthesize knowledge claims about a topic.¹⁰ A systematic review is like a literature review with one significant difference. A systematic review consists of a formal process with defined steps designed to capture every relevant piece of evidence and is the focus of the study itself. In contrast, a literature review is informal and targeted at specific concepts of interest to the researcher and is a means to an end, used as an introduction to the fundamental part of the research.¹¹

Systematic reviews typically take six to 18 months.¹² A typical systematic review process includes formulating the problem, searching the literature, screening for inclusion of relevant data, assessing the data quality, extracting the data, analyzing and synthesizing the data, and reporting the findings.¹³ Data in this study are extracted from academic journal articles, book chapters, white papers, magazine articles, etc. Because of the timeline constraints, this systematic review had to modify its approach by not including a formal quality appraisal assessment of articles considered. However, articles were informally assessed for quality. For example, articles were only included if authored by individuals who had credentials intersecting the IC and AI/ML.

Systematic Review Method

Key methodological aspects of a systematic review include the type of question asked and how the data is analyzed. Systematic reviews can ask different types of review questions. A descriptive type is used to support this study because it allows the researcher to characterize the state of knowledge of a topic, which is precisely what needs to be done in this study.

This systematic review used an interpretative approach.¹⁴ Interpretation involves a coding system to identify researcher-created categories and themes from the literature. Coding in this context should not be confused with computer science, where coding is a set of required instructions that make a computer do what computer scientists expect.

In an interpretive approach, a code is a word or set of terms generated by the researcher to capture the essential meaning of someone else's narrative selection, whether spoken or written.¹⁵ It is an interpretation process. The coding process involves several steps that begin with breaking down a narrative into component pieces, the analysis function. It ends with building a new framework from the pieces, the synthesis function.¹⁶

The coding process can be inductive and deductive. Inductive coding is data-driven and flexible, whereas deductive coding uses predetermined codes to organize data coding. Coding can use a combination of inductive and deductive processes.¹⁷ In this study, a combination of inductive and deductive coding is used.

Deductive coding serves as the initial framing device by using three predetermined components, each representing the review questions discussed above: motivation, problem, and solution. Authors who write about problems and solutions have their reason for why they are writing.¹⁸

Once the literature is separated into deductively generated concepts of motivation, problem, and solution, inductive coding begins and is data-driven. Within the motivation component, inductive coding is guided by the question: What is the motivation for writing about [topic] in the IC? Within the problem component, inductive coding is informed by the question: What problems are identified about [topic] in the IC? Within the solution component, inductive coding is led by the question: What solutions are recommended for [topic] in the IC?

Data Collection Strategy

Three search streams found 2,060 articles related to the IC and AI/ML. Each was reviewed for consideration with the question: What is the state of knowledge about problems and solutions within AI/ML in today's IC? Figure 2 summarizes the inclusion and exclusion flow from the three search streams to the final number of 41 articles used in this study.

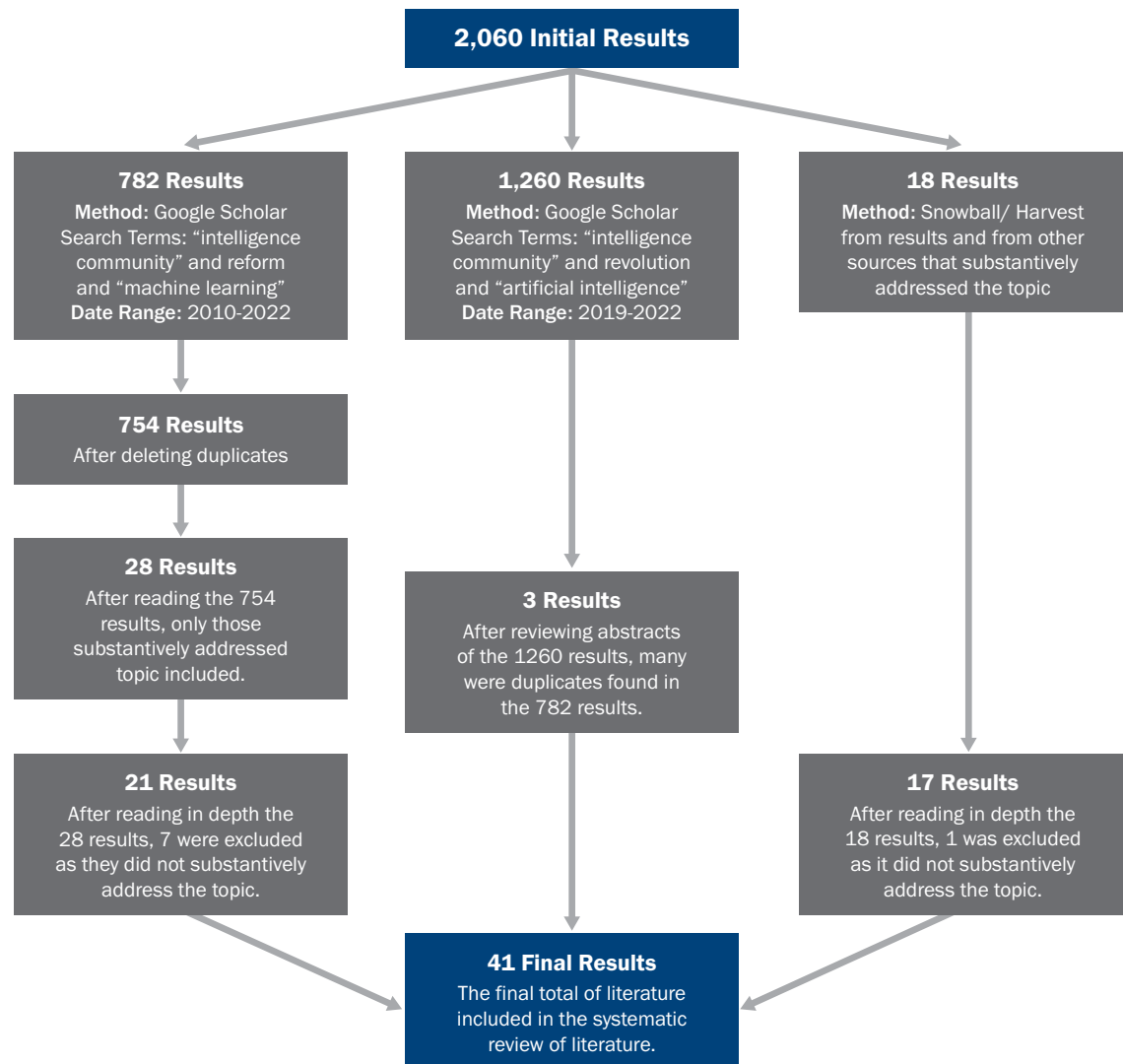
In the first stream, 782 results were retrieved from a Google Scholar search using the search terms “intelligence community,” AND reform, AND “machine learning” from 2010-2022, deleting citations. Of the 782 results, 28 were duplicates, resulting in 754 possible articles meeting the search criteria. After reading the 754 articles, it was determined that 28 pieces contained relevant information to answer the review question.

The articles excluded from the 754 were not focused on AI/ML or the IC. For instance, a mention of the IC and AI/ML might only be found in the introduction, recommendation, or reference sections where one sentence surmises that the research or findings would be of value to AI/ML and the IC. Similarly, the phrases “machine learning” and “intelligence community” were contained with reference titles. A second, closer reading of the 28 articles eliminated seven more because they failed to substantively address the review question. This left 21 articles to be included in the AI/ML review.

In the second stream, 1,260 articles were retrieved from a Google Scholar search using the terms “intelligence community AND revolution AND artificial intelligence” from 2019-2022. A shorter date range was used for the second stream of search terms than the first stream because a large number of articles were identified. Over 7,000 were retrieved if 2010 was used as the starting point. Almost all the second stream results were duplicates of the first stream, resulting in only three additional articles.

In the third stream, 18 articles were acquired through various non-systematic techniques, such as harvesting pieces from the first or second stream and tapping into various unclassified electronic communications about AI/ML, i.e., emails, web pages, etc., that included information of interest. After a detailed, closer reading of the 18, one was excluded due to its lack of substantive relevance to the review question. As Figure 3 shows, 41 articles were included to address the AI/ML topic within the IC.

Figure 3. Search Flow of AI/ML Systematic Review



Search Details

Most of the 41 articles included in the final review were published within the last two years (2020-2021). The search range for the first stream was 2010-2022, yet very few articles were published before 2019 related to AI/ML and the IC. The second stream search range of 2019-2021 included one from 2020 and two from 2021. The third-stream harvesting approach was mostly from 2021. The cut-off date for the 41 articles

is February 2022. Figure 4 provides the distribution of the final 41 articles used for this study by their publication year.

Figure 4. Distribution of Final 41 Articles by Publication Year

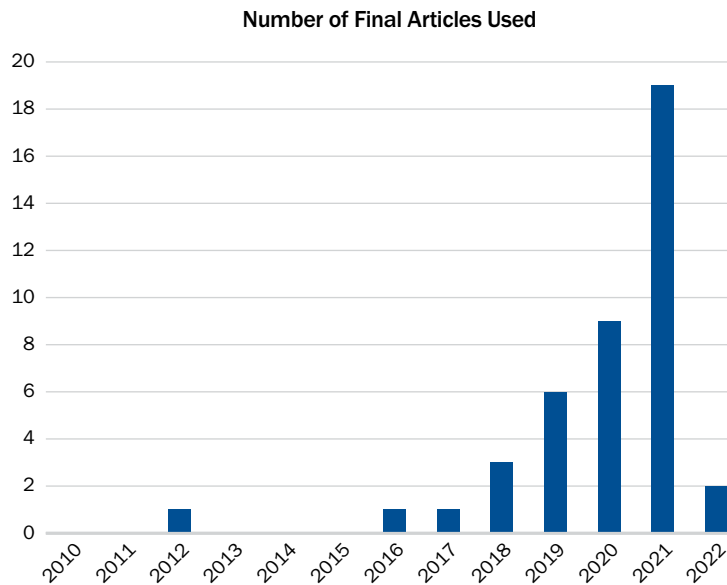
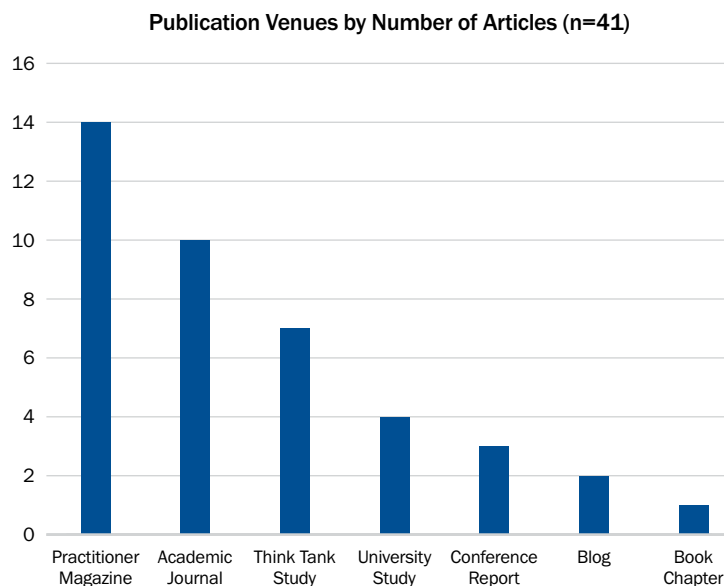


Figure 5. Types of Publication Venues

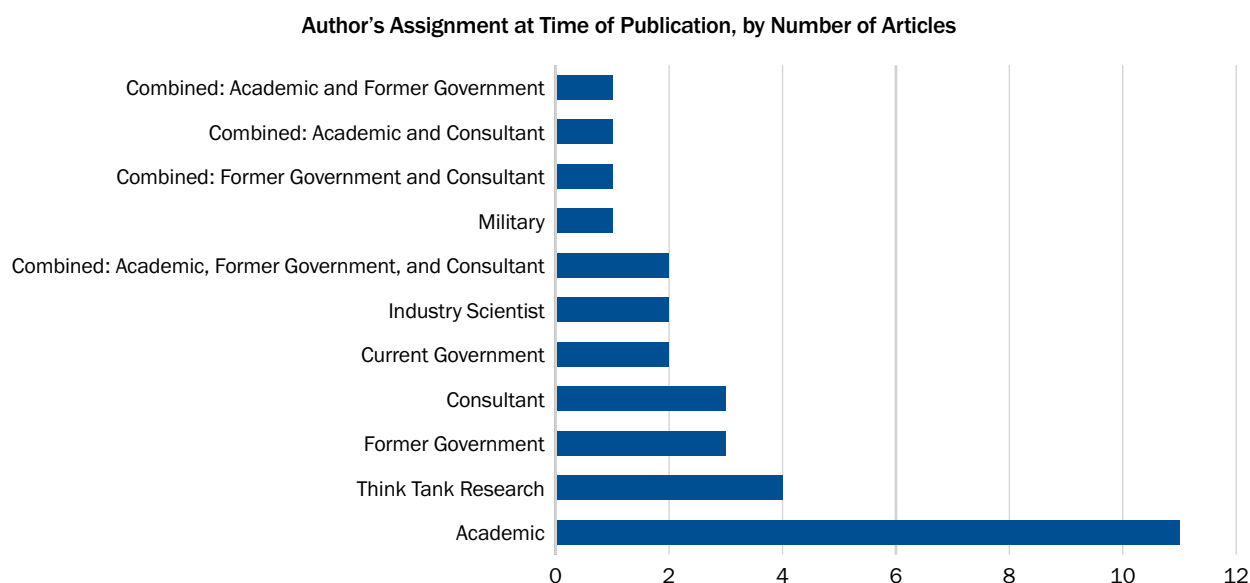


Articles were selected based on their relevance to the AI/ML overall review question. There was no restriction on the source of evidence, consistent with the methodological approach of a systematic review.¹⁹ As a result, a diverse selection of publication venues was considered reportable for the 41 AI/ML articles. Because a systematic review requires the identification of articles considered, a reference section is included at the end of this study. It contains the citations of all 41 articles and identifies the 41 articles with an asterisk in front of the author's name. Of the 41 articles, 14 were from practitioner magazines, 10 were from academic journals, seven were from think tank studies, four were from university studies, three were from conference materials, two were from Internet blogs, and one was from a book chapter. Figure 5 summarizes the distribution of publication venues.

Those who authored the 41 articles came from various professional backgrounds. Thirty-one authors wrote the 41 articles. A few authors wrote two articles, and one author wrote a series of five. Identifying their professional assignment at the time of publication provides context for understanding their knowledge

strengths and possible motivations for writing about topics. Not all individuals remain in the assignment they held when the article was published. For example, the current Director of National Intelligence, Avril Haines, was employed as a think tank researcher. Thus, her assignment at the time of publication was categorized as a think tank researcher, not a current or former government official. For the AI/ML topic, eleven were written by academics. Four were authored by think tank researchers and three were authored by former government individuals and consultants. Some articles were written by a combination of professionals. Figure 6 summarizes the author's job at the time of the article's publication.

Figure 6. Author's Job at Time of Publication



Analytic Strategy

Synthesizing the literature begins with the coding process. There are four increasingly abstract levels of inductive coding: codes, categories, themes, and concepts.²⁰ Codes are the interpretation of the claims provided by the authors. The claims are typically either embedded in a sentence or the entire sentence and selected by their relevance to the review question. Categories are the groupings of codes that share similar meanings. For example, a category about risk would represent the many codes and the evidence from source sentences in articles that discuss risk-related activities. Categories remain grounded close to the meaning of codes and can be thought of as a reflective process because categories reflect the similarities of the codes.

Themes, on the other hand, bring together disparate categories into a theoretically meaningful group and can be thought of as a formative process because they capture the differences between categories. For example, in addition to risk, other categories might include budgets, human resources, acquisition, etc. These

different categories form a strategy development theme because the categories involved the need to create a strategy. Concepts represented the most abstract way of capturing the meaning of themes. For example, strategy development was just one theme that spoke to implementation needs required before and during the deployment of an AI/ML system in the IC.

Table 1 provides the number of codes, categories, themes, and concepts for each review question: the motivation of authors who wrote about issues, the problems they identified, and the solutions they offered. The following section is organized by the review questions: motivation, problems, and solutions.

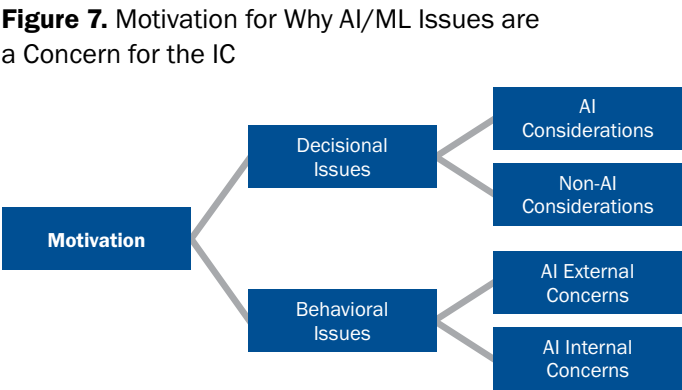
Table 1. Review Questions and Coding Results

Review Question (RQ)	Codes	Categories	Themes	Concepts	Total
RQ1: Motivation	51	13	4	2	70
RQ2: Problems	262	52	11	3	328
RQ3: Solutions	176	34	6	2	218
Total	489	99	21	7	616

Findings

Motivation

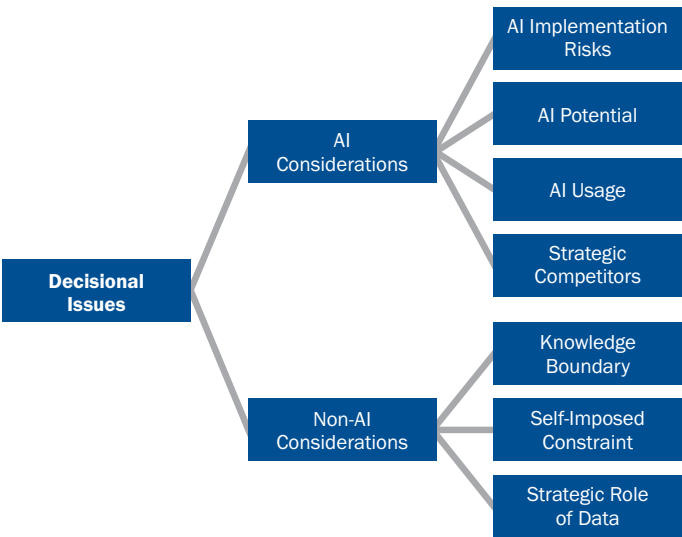
Motivations for why articles were written fall into two types of issues: those that relate to factors involved with decisionmaking about using AI/ML, called *decisional issues*, and those that relate to concerns about AI/ML technology behaviors, called *behavioral issues*. These two types of motivation refer to why an article was written about AI/ML in the IC. Figure 7, shown below, summarizes the two types of motivation concepts and their respective details.



Decisional Issues

Decisional issues consist of two types as shown in Figure 7: those that involve AI in decision considerations, and those that do not involve AI in decisions, non-AI considerations. *AI Considerations* include AI/ML implementation risks, AI/ML potential, AI/ML usage, and strategic competitors. *Non-AI considerations* include knowledge boundaries, self-imposed constraints, and the strategic role of data. The decisional issues are summarized in Figure 8 and discussed in detail below.

Figure 8. Motivational Decisional Issues



AI Considerations

AI implementation risks include concerns about the risks of using AI/ML involved in intelligence analysis, where deliberation is

typically the mainstay of critical thinking.²¹ The issue is whether analysts trust that AI/ML conclusions are accurate, because analysts directly interact with their organization's leadership and external customers.²² The challenges and dilemmas of using AI/ML in an intelligence context are caused by the well-known extent of the unknowns and ambiguities within the analytic profession.²³

AI potential means concerns about the anticipated increased use of AI/ML. These concerns include the need for a conceptual framework to help think about AI/ML and its relevance to the intelligence mission. Understanding a new phenomenon requires a theoretical underpinning to interpret the phenomenon.²⁴ Another concern is the need for metrics to assess the impact of AI/ML in support of intelligence missions, as measures of effectiveness are challenging when the end state of AI/ML use may not be conceptualized at this point.²⁵ In addition, an assessment of how technology assists to identify and prevent threats, an intelligence analyst's activity, has not been established.²⁶

AI usage addresses concerns about how AI/ML technology can be used. These include proposals for using AI/ML to avoid strategic surprise, as IC warning failures have resulted in significant surprises that have cast dispersion upon the IC.²⁷ Adopting AI/ML so that current intelligence processes can be improved, and understanding how AI/ML will be integrated into existing systems or replace existing systems, is still a work in progress.²⁸ Integrating AI/ML into efforts to deter threats, which involves knowledge beyond collection and analysis, includes broader instruments of power, such as diplomacy and economics, that interface with intelligence.²⁹ How to pay more attention to AI/ML testing and evaluation is essential so that analysts and consumers of intelligence will have more trust in their conclusions, as analysts and intelligence agencies will want to know to what degree their technologies are fully operational.³⁰ Finally, agencies will want to know to what degree their analyses are comparable with other agency products. There is a need to develop IC-wide standards or best practices.³¹

Strategic competitors concerns include great power's use or possible use of AI/ML against the United States and the asymmetric nature of AI/ML. There are signs of how great powers such as China and Russia may use this technology for nefarious, below-the-threshold-of-war activities.³² The re-emergence of near-peer competitors such as China and Russia are of concern. The need to adopt a focus on their strategic threat, which is a change in emphasis from the U.S. and allied focus on international terrorism over the past twenty years³³ as well as China's and Russia's allies, for instance, North Korea and Iran.³⁴ Lastly, the concern includes the need to preserve power primacy over China in light of China's persistence in preventing U.S. primacy.³⁵

Non-AI Considerations

Knowledge boundary issues include how the intelligence analyst's work at the task level may change. Understanding the boundaries between what AI/ML can provide and what the human analyst will do will be important.³⁶ The distinction between how vast quantities of information and intelligence are explained and understood is complicated because the edge of data as intelligence ends and intelligence from data begins may be clouded.³⁷ Lastly, the concern includes the degree that intelligence analysis adopts data science into its process and may introduce benefits or limitations. As the volume of data becomes ubiquitous, having a science of data, data science, becomes a necessity.³⁸

Self-imposed constraint issues include the lack of innovation in the IC culture. The culture reinforces the status quo rather than experimenting with something new.³⁹ Another issue is that the IC budget process is rigid, so it takes years to incorporate a new program.⁴⁰ Because the IC is risk-averse, and the adoption of AI/ML is new, risks must be taken, but adoption of AI/ML presents additional challenges without a risk framework.⁴¹ Lastly, the IC culture tends to throw money at problems, but adopting AI/ML will take more than just money.⁴²

The *strategic role of data* issues includes the idea that the Internet has no boundaries and is easy to access from anywhere and to almost everyone. Most of the world now has data democratization.⁴³ Further, because data is democratized, data must be conceptualized as a strategic asset; otherwise, adversaries and competitors will benefit from a fast response and the versatility in achieving a desired effect.⁴⁴ Acquiring new ways for the IC to leverage data collection, processing, and exploitation will be paramount.⁴⁵

Behavioral Issues

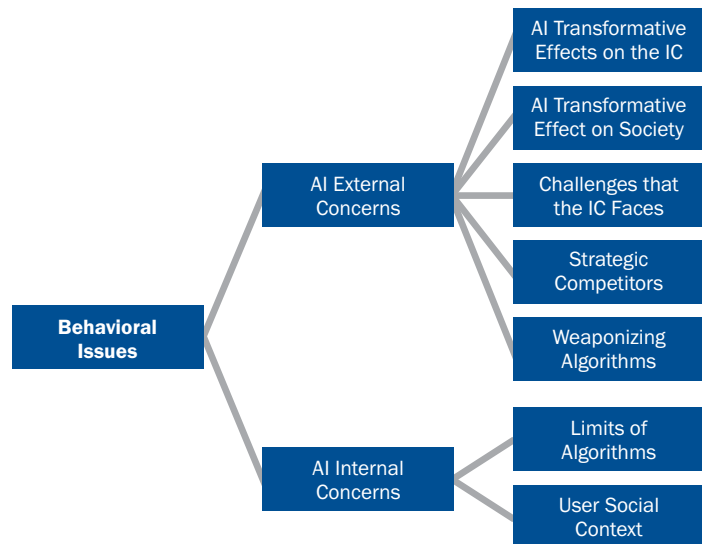
Behavioral issues consist of two types. *AI external concerns* originate outside the IC. *AI internal concerns* involve concerns from within the IC. AI external concerns include the transformative effects on the IC, the transformative effects on society, IC challenges with AI, strategic competitors, and weaponizing algorithms. *AI internal concerns* include the limits of algorithms and the AI user's social context. The behavioral issues are summarized in Figure 9 and discussed in detail below.

AI External Concerns

The concerns about the *AI transformative effects on the IC* include the degree to which analytic tradecraft is dependent on accessing large amounts of data, modeling the right architecture for the data, and ensuring the contextual understanding of the user is considered.⁴⁶ The human resource impact on intelligence analysis, the workflow, and the IC workforce is considered a significant challenge.⁴⁷ The system integration issues addressing the transformative effects of AI applications on existing and future system integration are not fully understood.⁴⁸

The broader concerns about the *AI transformative effect on society* include a wide swath of issues that will likely trickle down into intelligence agencies' concerns. These concerns include the degree to which AI/ML innovation will disrupt global political, economic, and social relationships.⁴⁹ They also have the degree to

Figure 9. Motivational Behavioral Issues



which the economic activity of AI/ML will change how businesses operate at the business-to-business and business-to-consumer levels.⁵⁰ The question of what transformative effects these breakthrough technologies will eventually have on international conflict remains an unanswered question.⁵¹ The pace of change from these converging AI/ML technologies and their impact on society is a work in progress.⁵² How the pace of change will change global threats and the IC's ability to detect them remains unknown.⁵³ The degree to which the United States is prepared to compete or thwart threats in an era of AI/ML has yet to be fully addressed.⁵⁴

The *challenges that the IC faces* with AI/ML include the IC not responding fast enough to leverage the innovative effects of AI/ML technologies.⁵⁵ The IC is not reacting fast enough to exploit the increasingly massive amounts of open-source data using advanced technologies.⁵⁶ The lack of details on how AI/ML and data science fit into the analytic workflow has negative implications for overall integration.⁵⁷

The concerns for *strategic competitors* and how adversaries are *weaponizing algorithms* in AI/ML include their use to shape the United States and its allies' hearts and minds through disinformation campaigns.⁵⁸ The role algorithms will have in future international conflict is especially concerning when linked with networks and sensor grids.⁵⁹ The ability of adversaries to attack the U.S. AI/ML systems is a strategic concern because they can deliberately introduce bias to increase an adversarial advantage.⁶⁰ If U.S. companies can acquire adversary AI/ML algorithms, then the U.S. companies can test their own algorithm's resilience against those of the adversaries or competitors, which would benefit the IC who procure AI/ML systems.⁶¹

AI Internal Concerns

The literature expresses numerous AI/ML concerns about the *limits of algorithms* within the IC. These limits include the general lack of understanding of AI/ML algorithms⁶² and the lack of AI/ML algorithmic capabilities in analytical tasks that require a degree of abstraction, such as intelligence planning, dissemination, and evaluation.⁶³ Another significant concern is the lack of understanding of where and how human bias enters into algorithm's construction. Bias can degrade the functionality and effectiveness of AI/ML technology.⁶⁴ Being unable to reapply an algorithm designed for a specific intelligence function with specific data to another context may be a counterintuitive limitation, but it is a real limitation of the current technology. Even if the data sets are similar, domain knowledge for each context is different.⁶⁵ The design limitations of algorithms, each created for a specific purpose, shape the design process and are not widely understood.⁶⁶

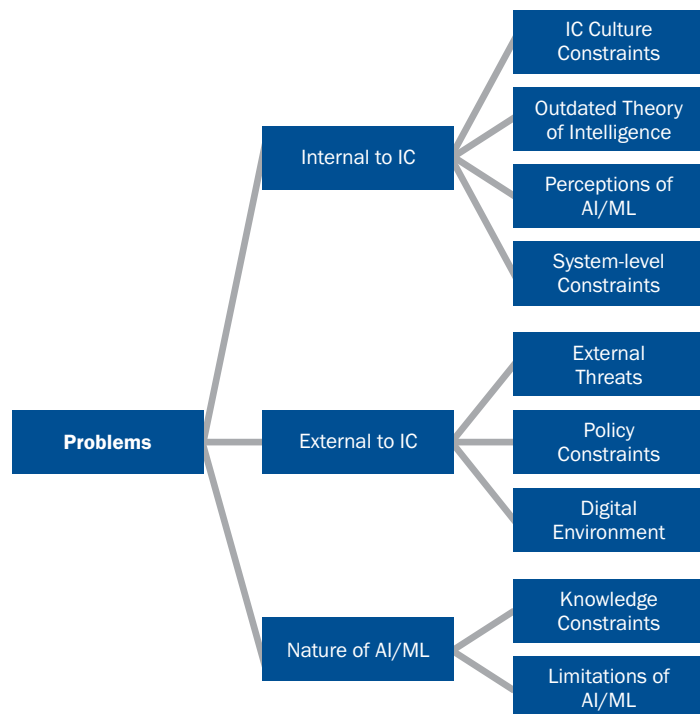
AI/ML technology exists within a *user social context*, which involves other analysts, senior analysts, and leadership. Within the analyst's social context, various articles raised several concerns, including the vital role of the analyst's interaction with analytical leadership.⁶⁷ These also include the analysts' perspectives on their interaction with AI/ML technologies⁶⁸ and how the social interaction between the analyst and the AI/ML technology affects the knowledge created for the decisionmaker.⁶⁹ Related, the interaction includes how analysts manage the social dynamics between organizational cultures within their own intelligence agency, other U.S. intelligence agencies, and external partners such as academia and industry.⁷⁰ The degree to which the AI/ML technology can anticipate the analyst's capabilities and intentions, especially in new situations, is unknown.⁷¹

Problems

Identifying problems with AI/ML as they pertain to the IC constitutes the bulk of the content from the articles used in this study. Problems are divided into three types: problems that the IC has direct control over, *internal to IC*; problems that the IC does not have direct control over, *external to IC*; and problems related to the AI/ML technology, *nature of AI/ML*.

Within each type of problem, more detailed perspectives are discussed. Figure 10 summarizes the three types of problems and their sub-problems, which are discussed below.

Figure 10. Problems of AI/ML for the IC



Internal to IC Problems

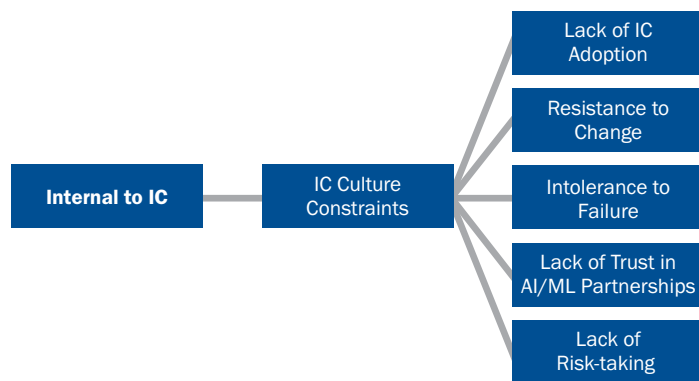
Internal to IC problems consist of four types. One type involves the IC culture's norms, customs, beliefs, values, and symbols or *IC culture constraints*. A second type involves what authors considered an outdated conceptualization of intelligence or *outdated theory of intelligence*. A third type involves the negative perceptions of AI/ML by analysts, leadership, and customers. A fourth type involves technology and data-related constraints of existing systems within IC agencies or *system-level constraints*.

Each of the four types is further divided into specific problems, which are discussed below in more detail. *IC culture constraints* are summarized in Figure 11 and include a lack of IC adoption, a lack of an IC digital shared vision, a lack of an innovative culture, and a lack of a risk-taking culture. The *outdated theory of intelligence*, summarized in Figure 12, includes challenges with big data, a Cold War approach, data incorrectly assumed as evidence, IC history of failures, lack of social science understanding, focusing on the macro-level threat, operating below the threshold of war, and having the simplified view that AI/ML is a solution to everything. In Figure 13, perceptions of AI/ML include negative analyst views, mixed leadership views, and negative customer views. Finally, *system-level constraints* are summarized in Figure 14 and include a lack of data standardization, a lack of AI/ML interaction testing, a lack of modern standards for testing, and a lack of system integration.

IC Culture Constraints

Cultural concerns are considered a fundamental problem for the IC. The *lack of IC adoption* is problematic since industry and academia have embraced AI/ML technologies and the data used.⁷² There is a

Figure 11. IC Culture Constraints Problems



perceived problem that the nation's competitive advantage could be lost within the next decade due to insufficient AI/ML adoption by the IC, further increasing the complexity of threats to the IC.⁷³

Resistance to change impedes an innovative IC culture and is one of the more important underlying constraints for AI/ML adoption. The primary obstacle to IC innovation is its culture⁷⁴ and the failure to think creatively about AI/ML across the various

intelligence domains.⁷⁵ The IC is behind the Department of Defense (DoD) in one aspect of confronting AI/ML because it has not developed the concept of pilot projects for the technology, has not received the same legislative tools, nor does it benefit from senior leadership support.⁷⁶ The IC is not proactive, and while voices have been heard with recommendations, rarely are they addressed until a catastrophe occurs.⁷⁷ The IC has not been able to adapt to the reality of data democratization, where the competitive advantage of classified information was once the sole purview of the IC. It is now threatened by open-source information.⁷⁸ Innovation has primarily been considered a competency of the private sector.⁷⁹ Intelligence agencies are not prone to promote an entrepreneurial spirit because of the lack of this competency in the IC.⁸⁰ Further, the lack of an entrepreneurial spirit is not universally welcome in the IC, and introducing new ideas is often discouraged.⁸¹

There is a *lack of tolerance for failure* in the IC, mainly because of the negative consequences of surprise. In this light, all failures, whether big or small, are considered equal.⁸² This means the U.S. Government does not like failures.⁸³ Consequently, the evolving nature of AI/ML algorithms means they will inject error and uncertainty, which requires a tolerance for failure.⁸⁴

A *lack of trust in AI/ML partnerships* is due to the absence of a shared IC digital vision, which is considered a central problem, especially since the technological expertise resides in the private and academic sectors.^{85,86} Challenges with implementation compound the partnership problem as it is one significant hurdle to attaining a shared vision across agencies. Still, it is quite another hurdle to implement a vision through standard competencies.

One of the main reasons for the lack of innovation is that the IC *lacks a risk-taking culture*.⁸⁷ There are many observations about this problem. There is no IC-wide mechanism for discussing and thinking about risks, taking risks, or offering perspectives on mitigating risks associated with AI/ML technologies.⁸⁸ There is a lack of an incentive structure to promote risk-taking, reinforcing the value of a normal business-as-usual mindset.⁸⁹ The existing IC culture, like that of the DoD, is focused on defending the United States. This implies that failure is not an option, further contributing to risk-taking resistance.⁹⁰ The IC culture likes repetition and routine, and its no-failure mindset, repetition, and routine provide a degree of well-being that analysts and managers are on the right path to ensuring security.⁹¹ While risk assessments are focused

on the risks of action, they often do not include the consequences of not taking action.^{92,93} The U.S. Government is generally risk-averse for the same reason as the IC. A failure to protect U.S. national security is not valued because the national security enterprise, including the IC, is paid for by taxpayer dollars.⁹⁴

Outdated Theory of Intelligence

AI/ML is often the technology that manipulates big data, while academia focuses on the three *challenges of big data*: volume, variety, and velocity. The IC has an additional set of challenges it must include in its theory of intelligence regarding big data. This includes veracity (accuracy and authenticity), volatility (latency), and value (noise is decreased).⁹⁵

Because of the volume of data, there is a mistaken belief in the IC that *data is incorrectly assumed as evidence*. While sometimes data is true evidence, finding the signals in the noise is not so apparent.⁹⁶ This mistaken belief that volumes of data assume a fact is in error because the significance of the data is not present within the data itself.⁹⁷

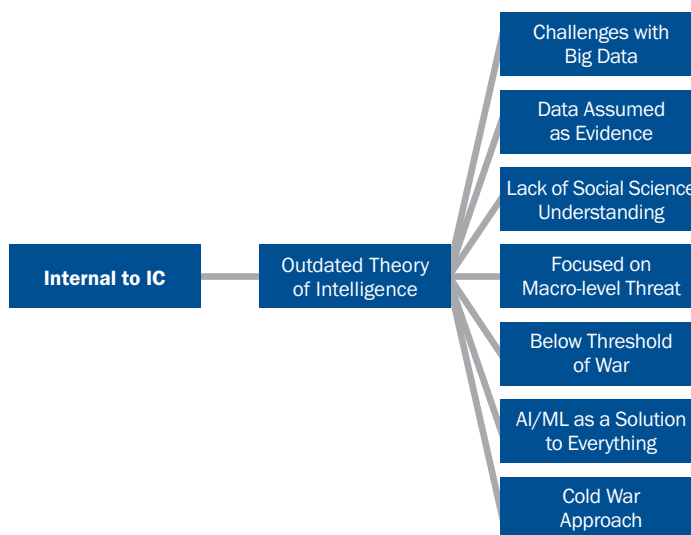
The IC suffers from a *lack of social science understanding*. One of the problems identified when seeking solutions through AI/ML technology is that technology takes a

front seat. Social science is underappreciated, especially as a science that enables the technologies.⁹⁸ A consequence of such an underappreciation is that the AI/ML algorithms identify correlations, which are not the same type of knowledge as causal models based on hypotheses.⁹⁹ Without a valid theoretical foundation, data analysis will likely lead to inferior technologies and less accurate answers.¹⁰⁰ The IC has not invested in the use of social science, which means the scenarios analysts construct cannot be effectively modeled or tested.¹⁰¹

The IC is *focused on the macro threat*, whereas knowledge of the local threat context is more important, yet difficult to assess. The IC invests in technologies that, at the regional-scale level, do not provide sufficient context at local, smaller scales necessary to detect the emergence of novel situations.¹⁰² Some of the questions that analysts confront are so complex that their answers reside in sub-populations or sub-entities outside the boundaries of what is normally accessible.¹⁰³

The reality of the 21st century, as of this time, is that *adversaries operate below the threshold of war*. China, Russia, and Iran show they are capable of doing so using a variety of techniques in an attempt to gain strategic leverage.¹⁰⁴ These techniques include using information communication technologies to target

Figure 12. Outdated Theory of Intelligence Problems



emotions and actions.¹⁰⁵ The activities aim to interrupt and interfere with all aspects of society, including efforts to reduce intelligence collection and pattern analysis.¹⁰⁶ These activities include softening populations to disrupt or change social norms, typically activities that intelligence has not been so adept at detecting.¹⁰⁷ The nature of these activities is not so black-and-white as to identify them as illicit or not; consequently, intelligence may not detect them.¹⁰⁸

As the promise of AI/ML technologies appears on the horizon for the IC, there is a tendency in the IC to *view AI/ML as a solution to everything*. This tendency manifests itself as the search for the holy grail, the magic button, that, consequently, has a negative impact on tradecraft.¹⁰⁹ The old adage applies, if you have a hammer, then everything must be a nail in order to use it. Big data is viewed in the same way. It is a magic button in place of the work that intelligence analysts pursue at the task level.¹¹⁰ While the magic mentality may exist, it cannot be taken for granted that AI/ML will decrease the need for human intelligence analysts.¹¹¹

According to the articles included in this study, the IC remains beholden to the 1945-1990 *Cold War approach to intelligence*. The legacy tradecraft relies on highly trained and experienced analysts who can identify the wrong data and the wrong process. Looking forward, analysts are already overwhelmed with data that, for the most part, cannot be identified as right or wrong.¹¹² As a result, analysts have a puzzle-based approach to intelligence where analysts just need to find the right pieces of intelligence to find the answer. This is an approach that is too narrow for today's information age.¹¹³ Additionally, the IC's history of failures has focused on the preventing strategic surprise, emphasizing prevention over anticipation.¹¹⁴ The IC focus and success on the tactical side of counterterrorism over the past twenty years will likely contribute to near-term strategic surprises.¹¹⁵

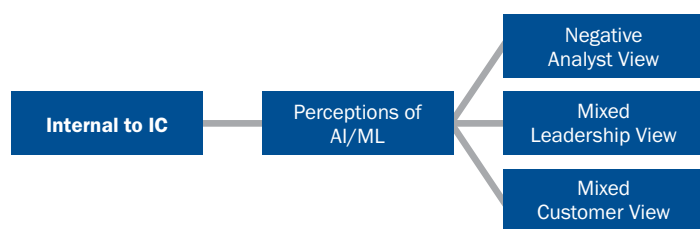
Perceptions

There is generally a *negative view towards AI/ML* by intelligence analysts for a variety of deep-seated reasons. Analysts are concerned that they will be held accountable for the mistakes made by AI/ML technologies that they forward to customers.¹¹⁶ The context and content surrounding an intelligence issue can influence

an analyst's thoughts, which is good. Still, it is not clear that AI/ML technologies can factor in such necessary considerations.¹¹⁷ Because the analyst's understanding of their own social and cognitive context is necessary for making sense of the data and conducting analysis, complexity is inherently part of the analytic equation.¹¹⁸ Analysts ignore or downplay unclassified data even

though they identify themselves as all-source analysts.¹¹⁹ Within the domain of unclassified data, open-source intelligence has largely been interpreted by intelligence analysts, in general, as press reporting,¹²⁰ which generates a perception that open-source information is less valuable than classified data.¹²¹

Figure 13. Perceptions of AI/ML Problems



In other analyst perception-related issues, AI/ML systems are procured and delivered for intelligence analysts to use without their involvement. Not because technology proponents have a disregard for analysts but because of the well-meaning assumptions of such proponents.¹²² Because of the lack of their participation, the intelligence analyst has not been able to share the challenges they face as analysts with technology proponents. Hence, analysts may view AI/ML technology solutions as not relevant or not helpful.¹²³ This lack of involvement negatively affects the analyst's level of trust in data and algorithms.¹²⁴ The tendency towards distrust about AI/ML is greater than any other human-related interaction with technical staff, management, or executives.¹²⁵

Consequently, the distrust has led to a lack of demand signals for AI/ML by intelligence analysts.¹²⁶ When distrust is compounded with a lack of user buy-in of AI/ML, intelligence analysts reinforce the reluctance to embrace AI/ML.¹²⁷ This leads to questions and concerns beyond trust by intelligence analysts that include the challenges of explaining the logic of the analysis, the robustness of the analysis, and the effectiveness of AI/ML.¹²⁸ These challenges open up a broader set of negative perceptions based on the analyst's lack of understanding the process by which AI/ML generates conclusions.¹²⁹ It seems common sense to require analysts to know about AI/ML technologies to apply them and integrate their output within the analyst workflow.¹³⁰ If not understood, this lack of understanding leads to analysts perceiving AI/ML as a black box, making AI/ML decisions difficult to understand and explain.¹³¹

Perceptions are mixed towards AI/ML by analytic leadership, though reported as more negative than positive. One problem is the dichotomy between leadership and intelligence analysts about technology. Decisions made by leaders do not consider analyst skills or expertise with AI/ML, suggesting that leaders assume their positive view of the technology is all that needs to be considered for success.¹³² This fallacy of perceptions is a problem because analysts will see the problem differently than leaders or customers. Leaders want a solution for dealing with a large amount of data and want computer scientists to develop the technology. In contrast, intelligence analysts who work on the problem bring their preferences and context.¹³³ This disconnect results in further misalignment between aspiration and implementation. Leadership says how important AI/ML is, but initial small or pilot efforts using AI/ML have a hard time scaling up to institutionalization. Scaling up is problematic because of the lack of attention by leaders during the complex scale-up attempt and the lack of a demand signal from intelligence analysts who do not see any urgency.¹³⁴

Other leadership problems with the perception of AI/ML include the dilemma, of the observation by leaders that the IC does not have its own AI/ML talent. On the other hand, the talent resides elsewhere. For one, the government is not a significant customer of AI/ML; as a result, most AI/ML researchers and scientists do not work for the government.¹³⁵ The private sector is cornering the talent marketplace.¹³⁶ This makes it difficult for the DoD and the IC to attract AI/ML talent to test and evaluate these technologies.¹³⁷

Consequently, private sector organizations produce intelligence that rivals the IC. It is faster and cheaper due to the integration of AI/ML. This is likely what IC leadership sees but may not understand the reason the IC cannot duplicate what the private sector is doing.¹³⁸ The IC wants commercially available AI/ML but then adds their specific requirements that offset cost savings and may limit the functionality of the commercially available AI/ML.¹³⁹

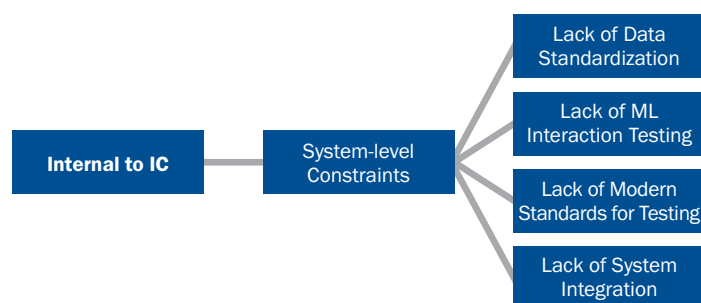
While IC leadership wants AI/ML, integrating the technology into the analyst's task-level workflow will require an extensive retooling of analyst competencies, skills, and knowledge.¹⁴⁰ Even if analysts develop such rethinking, the turnover of analysts with such new knowledge will become a second-order effect to solve.¹⁴¹ Such an effect is indicated by leadership's lack of understanding of the workings of AI/ML algorithms.¹⁴² It is not only IC leadership's lack of understanding but also policymakers who lack understanding.¹⁴³ This lack of understanding may result in erroneous decisions or conclusions.¹⁴⁴

There is also a *mixed customer view*. Customers are biased towards unclassified information and implicitly view the benefits of technology-assisted processing. This is especially true as intelligence agencies and non-intelligence organizations increase their reliance on open-source information and cheaply use AI/ML technologies, which can quickly process massive amounts of data.¹⁴⁵ Because when they do get classified information, it is too slow to arrive and too highly classified to be of use to customers.¹⁴⁶ However, customers are reluctant to trust that AI/ML products can be customized to meet their needs, which is compounded by the general lack of knowledge by intelligence analysts of the informational needs of policy customers.¹⁴⁷ Policymakers typically have more general, less technical questions for intelligence analysts. Still, it is in the details that the more general questions can be answered. Consequently, IC customers are reluctant to trust AI/ML because of their concern that AI/ML will not be able to support the synthesis from detail into the more general.¹⁴⁸ The IC fears that customers will have less of a demand signal for IC products and increased demand for private sector-created intelligence.¹⁴⁹

System-Level Constraints

The IC has a *lack of data standardization* across agencies. Data exist in many formats across systems that are either disconnected from each other or not accessible to each other. No standardized science defines the organization, management, and storage format. There is no standardized way of representing, naming, and categorizing properties and relationships.¹⁵⁰ Such disconnections make it hard to conduct data fusion and create all-source intelligence products.¹⁵¹ Because of the diversity of data and systems, analysts have to spend time figuring out what data is related to each other so that it can be pieced together across data sets and systems.¹⁵² In addition, data is inherently messy. This raises concerns about how AI/ML can be used to handle data disorder, which is a significant obstacle to using AI/ML tools.¹⁵³

Figure 14. System-Level Constraints Problems



There is a *lack of AI/ML interaction testing* between AI/ML technologies and existing or non-AI/ML technologies for which data would move between. There is a fundamental challenge of integrating AI/ML into existing technologies that involve understanding the costs, risks, and benefits of investing time and energy.¹⁵⁴ Unexpected failures can occur when AI/ML interactions are not tested.¹⁵⁵ The problem

is compounded by the difficulty in figuring out why an AI/ML system made a decision in a specific situation.¹⁵⁶

There is a *lack of modern standards for testing* AI/ML technologies. Current testing standards and methods are not optimized for AI/ML.¹⁵⁷ Neither are policies and metrics for testing performance and evaluating risk.¹⁵⁸ While metrics may be available for consideration during the development of AI/ML technologies, standards for operational performance have not been established.¹⁵⁹ There is a lack of iterative and continuous approaches to testing once the AI/ML has been acquired and undergoes evaluation before use.¹⁶⁰ This lack contributes to the difficulty in reevaluating AI/ML systems every time a system is upgraded, above and beyond the cost and effort of doing so.¹⁶¹

A *lack of system integration* exists as there are technical barriers with integrating unclassified data on the unclassified system and classified data on the classified system, as well as integrating these movements into an analytic workflow supported by AI/ML.¹⁶² This situation is further compounded by the lack of compatibility between AI/ML and existing systems, as ways of processing data or delivering products may differ.¹⁶³

External to IC Problems

External to IC problems consist of three types: threats from foreign actors (*external threats*, see Figure 15), policy constraints from the U.S. Government outside of the IC (*policy constraints*, see Figure 16), and the very nature of digital information (*digital environment*, see Figure 16). Each of the three types is further divided into specific problems, which are discussed below in more detail: *external threats* include the adversary weaponizing AI/ML, the underprioritizing adversary use of AI/ML against the United States, adversary easy access to AI/ML, adversary AI/ML use for deception, and adversary AI/ML use for disinformation. *Policy constraints* include a lack of appropriate oversight, a lack of flexible acquisition rules, and a lack of flexible budget rules. The *digital environment* includes information overload of digital devices, information overload of data from these devices, and the effects of information overload on analysts.

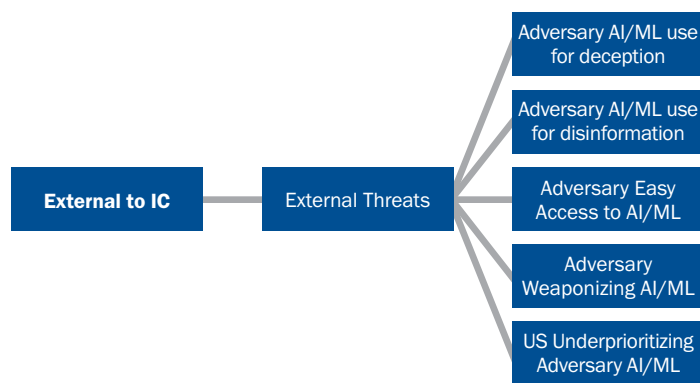
External Threats

The articles included in this study address concerns about foreign actors' use of AI/ML to create disinformation and deception. While the articles in this study do not define these two terms, they will be described here. Disinformation and deception are related, but they are distinct phenomena. Disinformation is information, and the information is designed to intentionally get others to advance some political or social end state by changing a belief system.¹⁶⁴ Deception is the outcome of disinformation due to an attitudinal or behavioral change caused by the actor's intention to mislead.¹⁶⁵

One of the problems is that *adversary AI/ML use creates deception*. Data is altered or manipulated by an adversary representing new opportunities for adversaries.¹⁶⁶ Russia's efforts have produced deepfakes, videos, and audio through AI/ML.¹⁶⁷ Adversarial use of AI/ML is a threat, especially to U.S. AI/ML algorithms.¹⁶⁸ These deepfakes can flood the IC data collection efforts to create undesirable outcomes.¹⁶⁹ Adversaries will use AI/

ML to limit U.S. data, thereby fooling algorithms.¹⁷⁰ As a result, the IC will be more vulnerable to deception, sources and methods exposures, information operations, cyber operations, and counterintelligence operations.¹⁷¹ Adversaries will also use AI/ML defensively to complicate, disrupt, and degrade IC efforts to collect

Figure 15. External Threats Problems



against adversaries.¹⁷² These adversarial AI/ML efforts collectively will have the effect of degrading American trust in its institutions as more deception takes place and adversaries find it easy to do.¹⁷³

Adversary AI/ML use creates disinformation is also a problem.¹⁷⁴ Adversary AI/ML-biased algorithms are used to create disinformation to advance adversary agendas, thereby undermining the U.S. or allied legitimacy.¹⁷⁵ Adversaries do this by taking advantage of the brittle nature of AI/ML algorithms, i.e.,

one small change and everything falls apart.¹⁷⁶ Adversaries are focusing on U.S. data collection efforts that mine social media. Injecting disinformation has the potential to introduce false information into analysis.¹⁷⁷ This will make determining the truthfulness and value of data collected much harder for the IC, especially with the volume of information collected.¹⁷⁸ Countering such adversarial use of AI/ML will require deploying a robust U.S. AI/ML technology.¹⁷⁹ The undesirable effects of hostile disinformation through social media have the potential to influence public opinion and cause panic on national security issues.¹⁸⁰

Adversary easy access to AI/ML is ubiquitous.¹⁸¹ Because of the pace of change, the adversarial impact of AI/ML has a significant negative effect on the strategic environment.¹⁸²

Adversarial weaponizing of AI/ML is the underlying strategy that produces disinformation and deception. The early indicators of such weaponization are clear, and the IC needs to address them.¹⁸³ It is not only U.S. military infrastructure that is at risk from such weaponization, but also the civilian infrastructure.¹⁸⁴ AI/ML provides adversaries with entirely new capabilities outside traditional military capabilities.¹⁸⁵ One counterintuitive phenomenon is that the United States addresses adversarial weaponization in the public forum. This allows adversaries to circumvent U.S. detection, capture, and defeat capabilities.¹⁸⁶ One way this can be done is by adversaries, or the private sector, creating biased data that could serve as bait with the hope that data collectors, like the IC, who use the data for training algorithms purposes, would degrade the quality of the U.S. AI/ML algorithms.¹⁸⁷

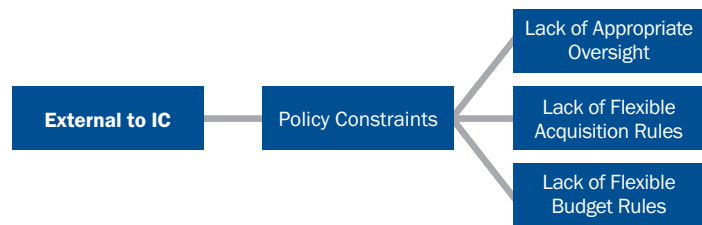
Adversaries are now aggressively pursuing AI/ML technology to use against open-source information, not just social media, and, so far, they are more capable than the IC.¹⁸⁸ Adversaries are harvesting data on American individuals and building profiles to manipulate or coerce them.¹⁸⁹ A wide array of adversaries is pursuing these capabilities.¹⁹⁰ One consequence is a counterintelligence threat. Case officers will struggle to maintain cover with risk to themselves, their agents, and operations.¹⁹¹ If adversaries gain access to the control information infrastructure, they can control the data, which can then be compromised.¹⁹²

The articles note that the United States has *underprioritized adversary use of AI/ML*. This observation falls mainly in the hands of decisionmakers who are biased against developing U.S. AI/ML capabilities.¹⁹³ Adversaries will use AI/ML as part of their military capabilities to improve their weapon performance and counter U.S. and allied weapon systems.¹⁹⁴

Policy Constraints

There is a *lack of appropriate oversight* regarding the IC's use of AI/ML. Congress is risk-averse, which is an obstacle to the IC acquiring and using AI/ML.¹⁹⁵ Compounding this reluctance is that IC committees are currently organized along IC agencies.¹⁹⁶ Congressional committees are focused on certainty, clarity, and stability, which are expectations that directly oppose the fluidity of emerging AI/ML technologies.¹⁹⁷ The oversight provided by the committees to the IC, House Permanent Select Committee on Intelligence, Senate Select Committee on Intelligence, House Committee on Appropriations-Defense, and Senate Committee on Appropriations-Defense, has not evolved with the advent of the information age.¹⁹⁸

Figure 16. Policy Constraints Problems



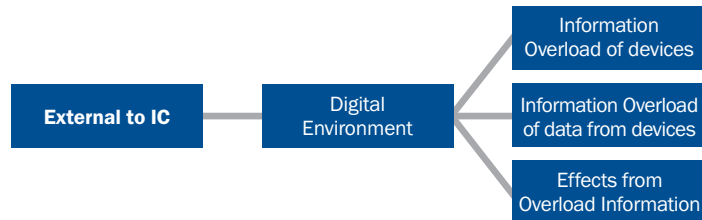
On the procurement side, there is a *lack of flexible acquisition rules*. Acquisition policies thwart a quick response capability and, in general, are too strict, with laws and regulations that are biased towards the accuracy and requirements of specifications instead of achieving the desired outcome.¹⁹⁹ As a result, the current acquisition process and cultural norms designed for physical systems do not support AI/ML technologies.²⁰⁰ These outdated practices thwart the IC's ability to adapt, restructure tasks, and modify AI/ML algorithms as needed because of the constant changes within the operational environment.²⁰¹

On the financial side, there is a *lack of flexible budget rules*. There is a dual concern: first, the rules of the budget process are so rigid that the ability to react quickly or quickly integrate AI/ML is severely limited; and second, the rapid pace in the development of AI/ML changes faster than the rigid budget system allows a response.²⁰² As a result, the budget process, which is so complex, inhibits innovation, which is needed to promote AI/ML technologies.²⁰³

Digital Environment

An *information overload of devices* handles digital data, especially smartphones.²⁰⁴ In 2020, the number of digital devices connected to the Internet was estimated to range from 30 to 50 billion.²⁰⁵ Because these technologies are so widespread, they generate massive amounts of information.²⁰⁶ This digital epidemic has increased the pace of life, yielding the greatest number of interactions between individuals in the history of humankind.²⁰⁷ These technologies produce structured and unstructured data that overwhelm analytic tradecraft and

Figure 17. Digital Environment Problems



pattern recognition capabilities.²⁰⁸ These devices have become the battlespace shaping adversarial engagement across such tools as Twitter, Facebook, Instagram, etc.²⁰⁹

The *information overload of data* includes data that is available commercially and publicly, which are open to almost anyone in the world having Internet access.²¹⁰ The speed

and volume of this data democratization overload all intelligence domains except possibly human intelligence.²¹¹ The problems of managing and differentiating these data will intensify.²¹² Since IC does not know which data are important, all data must be collected and processed.²¹³ IC agencies are suffocating in data overload.²¹⁴ The volume of data collected, classified and open-source, greatly surpasses the analyst's ability to read and synthesize it for insights and relevant information for decisionmakers.²¹⁵ To quantify this overload, there are 40 times more bytes of data than there are stars in the observable universe.²¹⁶ This data extends far beyond what is in news feeds, whether print, Internet, or broadcast.²¹⁷

The *effects of information overload* on analysts include analysts becoming less confident in AI/ML judgments,²¹⁸ or by implicitly trusting AI/ML without critical thinking, they could become overconfident.²¹⁹ Within the collection domain, the overload makes prioritization challenging because of the dramatic increase, diversity, and shifting of targets and threats.²²⁰

Nature of AI/ML

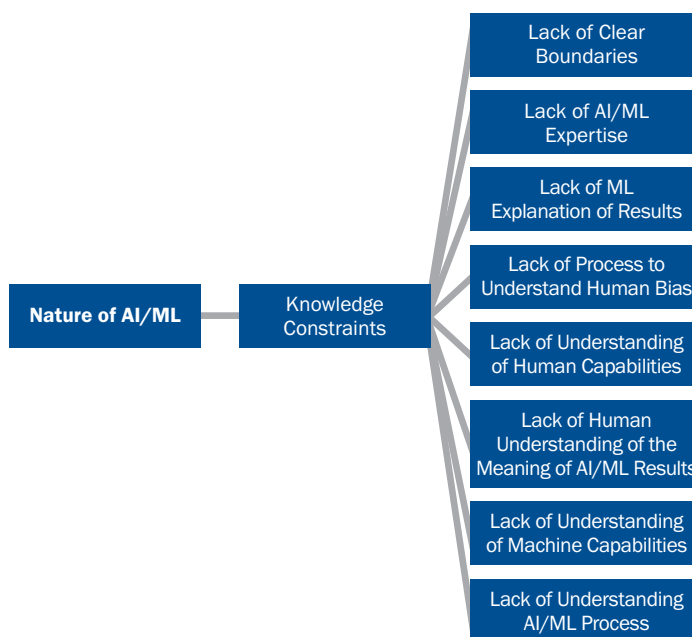
The *nature of AI/ML* problems consists of two types: the lack of knowledge by IC agencies (*knowledge constraints*, see Figure 18) and the technological limitations of the AI/ML technologies (*limitations of AI/ML*, see Figure 19). Each type is further divided into specific problems, which are discussed below in more detail. *Knowledge constraints* include a lack of clear boundaries, lack of IC AI/ML expertise, lack of AI/ML explanation of results, lack of process to understand human bias, lack of understanding the human capabilities by AI/ML, lack of understanding the machine capabilities by AI/ML, and lack of human understanding of AI/ML results and its meaning. *Limitations of AI/ML* include challenges with getting relevant training data, limited AI/ML algorithm design constraints, humans import bias into AI/ML design, lack of algorithmic transparency, and AI/ML not being good with deliberative tasks.

Knowledge Constraints

There is a *lack of clear boundaries* between various organizational-related entities. There is an ambiguity between specific and general thinking about what AI/ML can do for the IC. Both policymakers and analysts tend to think in the latter way when they should think in the former.²²¹ IC agencies and the IC tasks around the structure of intelligence domains, i.e., INTs, when such conceptualization is not relevant to AI/ML

technologies.²²² The fact that AI/ML technologies are blurring boundaries around data types means that secrecy is a greater risk than the old assumption that secrecy is highly valued because it reduces risk and uncertainty.²²³ Not communicating how an organization uses AI/ML to monitor behavior makes it almost impossible to determine how AI/ML is used to thwart adversaries.²²⁴ The increased connectivity worldwide makes it very difficult for agencies to determine who does what and where.²²⁵ In terms of research on AI/ML phenomena, academia does not have an interdisciplinary approach. Instead, each discipline, like cognitive psychology, robotics, neuroscience, etc., stands alone, which makes solving human-AI/ML teaming problems difficult.²²⁶

Figure 18. Problems Associated with Knowledge Constraints



The *lack of IC AI/ML expertise* is well recognized.²²⁷ AI/ML technologies require a new set of skill sets and knowledge from traditional analysis unassisted by such technologies. These include measuring, judging, and factoring in new attributes such as data authenticity.²²⁸

One of the big problems IC members experience is the *lack of AI/ML explanation of results*. This is compounded by each agency having its own cultural and institutional preferences for how and what an explanation of results should include.²²⁹ Analysts are hard-pressed to put faith in some results they cannot explain and defend at the analyst level.²³⁰ Analysts are charged with defending decisions to customers and leaders, and the limited functioning of AI/ML in this regard is not tolerated.²³¹ On the other end, there are concerns that AI/ML results may be so complex that the AI/ML result may not be explainable.²³² Overall, there are analyst concerns that AI/ML may only be viewed as a black box, a perception that the inner workings of a process are not knowable that may never be explainable to analysts.²³³ It is the complexity of the algorithm that is viewed as the black box. While the algorithm developer understands the algorithmic process, it is unclear whether the developer can transfer that knowledge to analysts.²³⁴ These concerns are not theoretical, whether developer or analyst, because if AI/ML results are not explainable, human judgment errors may cost lives.²³⁵

There is a *lack of process to help understand bias*. As noted above, human bias enters the AI/ML realm, intentional or not. However, the impact of bias by AI/ML technologies is not understood.²³⁶ There is currently no mechanism for analysts to collaborate to discuss ways of detecting or mitigating bias about AI/ML.²³⁷

There is a *lack of understanding of human capabilities*. There are challenges in understanding how the individual achieves common ground with a communication partner such as an AI/ML technology.²³⁸ It is unknown how individuals learn from a single event that has not occurred previously, as this would be

important in understanding an emerging, novel situation. Such a lack of knowledge may encumber the interaction of AI/ML technologies.²³⁹ There are challenges in understanding how individuals absorb the meaning of new situations and then are able to make predictions or generalizations.²⁴⁰ The challenges include understanding how individuals create mental models of situations, and other people's goals, intentions, and abilities, which would impact how the analyst would understand results from AI/ML.²⁴¹ In terms of understanding the machine-level terminology used to express its results, it is unknown what is required for individuals to gain the machine's trust.²⁴² Finally, how individuals perceive and understand the machine's mental abilities is unknown.²⁴³

There is a *lack of understanding of the meaning of AI/ML results*. Big data results can be faked and hacked because probabilities of correlations can be gamed.²⁴⁴ Big data is only an analytic tool.²⁴⁵ Big data generates correlations, but an analyst can err in judgment if the domain expert is not involved.²⁴⁶ Deriving erroneous meaning from results can occur from poor operationalization or searching for relationships that do not exist.²⁴⁷ In general, the opaqueness of AI/ML applications drives the inability of analysts to understand how the results came to be.²⁴⁸

There is a *lack of understanding of the machine's capabilities*. There are challenges in understanding what a machine needs to perceive, communicate, model, problem-solve, and learn with a human.²⁴⁹ Similarly, it is unknown which aspects and to which degree machines need to model the minds of individuals given the tasks at hand.²⁵⁰ There are challenges in understanding how machines construct a model of individuals to provide effective communication.²⁵¹

There is a *lack of understanding AI/ML process*. The ability of AI/ML to explain itself is problematic as analysts will want to know the logic, assumptions, and data biases of the algorithms used to create results.²⁵² Individuals do not know how AI/ML algorithms make the decisions they do.²⁵³ There is a mismatch between the priorities of the AI/ML and the metrics used, which makes it difficult to assess the system's performance.²⁵⁴ The effectiveness of AI/ML depends not only on the system's properties but also on how the system is used by analysts, making derivation of metrics challenging to develop.²⁵⁵ Metrics that determine if the AI/ML can duplicate human performance impede the value of the AI/ML.²⁵⁶

Limitations of AI/ML

There are *challenges with getting relevant training data* for AI/ML training. Getting large volumes of commercial data to test and train algorithms is difficult.²⁵⁷ Collecting and labeling data used for training is a challenge, especially for classified data.²⁵⁸ AI/ML using supervised training, which is what most systems use, requires large and representative data sets, and acquiring information is not easy.²⁵⁹ The tough challenge is that training data is often not real-world data, which means that AI/ML algorithm's interpretation from non-real-world data may not be useful in reality.²⁶⁰ Existing AI/ML applications are often flashy but not geared to the environment where classified data is integrated with unclassified data, thereby subject to poor judgments.²⁶¹ AI/ML systems must be trained not to collect or analyze unauthorized persons or entities, such as U.S. persons.²⁶²

It is an inherent problem that *humans import bias into AI/ML* in many forms. One problem is that if cognitive processing for analysts is offloaded to AI/ML, analytic decisions that form assessments will increasingly become less critical.²⁶³ The AI/ML algorithm developer cannot see their biases, which can manifest as an unexplainable error by the AI/ML.²⁶⁴ Because bias is so systemic, intelligence agencies and corporations that collect large data sets will likely introduce bias.²⁶⁵ Algorithmic bias is intro-

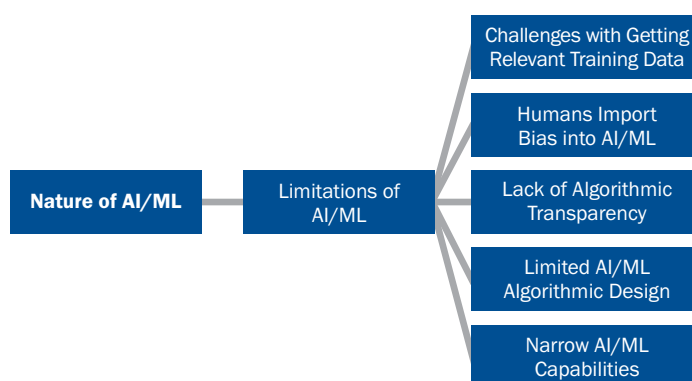
duced by whoever designs or trains the AI/ML.²⁶⁶ Bias can also be introduced during the acquisition process, from planning to deployment.²⁶⁷ Data scientists introduce bias.²⁶⁸ An interesting bias is that while humans are loss averse, AI/ML algorithms are not, yet humans design algorithms.²⁶⁹ The algorithm's perspective reflects the designer's organization, education, or experience in making design choices.²⁷⁰ Customers may introduce confirmation bias through their request for customization to address their needs and requests for certain algorithms.²⁷¹ In general, the pervasiveness of bias is constituted by many factors, including the characteristics of the intelligence data, the intelligence analyst, the programmer, the developer, and the data scientist.²⁷²

There is a *lack of algorithmic transparency* for analysts and customers. Analysts are concerned about how AI/ML arrives at conclusions.²⁷³ Algorithmic visibility is not very visible, especially as organizations incorporate more tools, and AI/ML internal processes, the black box, may not receive the attention it deserves.²⁷⁴ A confounding issue is the trust issue organizations assign to AI/ML, as the algorithms may give access to data that not all analysts have access to. This raises the question of how much trust can be placed in machines and how much trust analysts have in the algorithms.²⁷⁵ Access issues may thwart an organization's auditing for vulnerabilities that may be internally created by the organization or externally manifested by adversaries.²⁷⁶

There are many problems identified related to *limited AI/ML algorithm design*. There is the problem of space delineation, where algorithm designers have to understand the intelligence problem and transform the problem into a problem space. The level of transformation effort varies with the problem.²⁷⁷ Algorithm designers have to find ways for the algorithm to identify the intelligence topics, in which the identification involves processing a word or combination of words that are given weights.²⁷⁸ Another issue is the ability to match algorithms' cultural and institutional preferences with the working-level analyst's tradecraft standards.²⁷⁹ The brittle nature of algorithms is a systemic problem as the algorithm will only work as designed if the input data or the environment does not change.²⁸⁰ Big data is less robust when it is complex, which means algorithms using common data sets will be more successful than complex data sets.²⁸¹ However, most interesting data sets are dynamic, which change over time, meaning that the algorithm's correlations are only accurate for a specific period.²⁸²

An algorithm trained for an intelligence problem and scope will not be trained and usable for different intelligence problems or changes in scope.²⁸³ Algorithms cannot do deliberative thinking tasks such as

Figure 19. Problems Associated with Limitations of AI/ML



identifying what data to collect, who or how to disseminate the results, and assessing the risks and rewards of using results.²⁸⁴ The underlying challenge behind these limitations of algorithms is that their techniques are designed to quantify uncertainty or ambiguity.²⁸⁵ Algorithms do not work well when they lack regularity of context. The environment, statistical regularity, and populations are different during the data training phase and the training algorithm's deployment.²⁸⁶ These limitations make it difficult to figure out the actual capabilities and limitations of algorithms.²⁸⁷ This difficulty becomes more acute as algorithms become more complex because algorithmic output will be less transparent and less likely to predict the limits of its operation.²⁸⁸

The inflexibility of algorithms is further attenuated when faced with rare or novel events, as algorithmic design models cannot identify the meaning of such events.²⁸⁹ There is a relatively limited number of algorithm designs for which designers may select, which both introduce bias and incorporate unstated assumptions into the chosen design.²⁹⁰ This limitation is apparent when the need exists to recognize objects in different environments and under different constraints and therefore require different design solutions. For a hypothetical example, while an analyst may view two tasks as very similar such as recognizing a human face in a photo posted on an Internet web site and recognizing a missile launcher in a satellite image, the algorithm designer views these tasks completely different.²⁹¹

Related to the problem of a limited AI/ML algorithm design is the *narrow AI/ML capabilities*. AI/ML provides value when the input data can be unambiguously associated with output but is limited when the process is not so easy to map the input to output.²⁹² As such, it is better to think of the value of AI/ML as contributing to foundational intelligence, information about foreign military equipment, for example, instead of finished analysis where assessments are made about prediction, intention, or motivation.²⁹³ Related, AI/ML technologies are better at quantification and large volumes of data. Still, the intelligence analyst is better than AI/ML when the intelligence problem involves nuances of context and understanding of social interactions between actors.²⁹⁴ Hence, AI/ML technologies are limited in their capacity to deliberate.²⁹⁵ AI/ML technology will not provide immediate answers to any problem space having dynamics, variation, and temporal complexity.²⁹⁶

Solutions

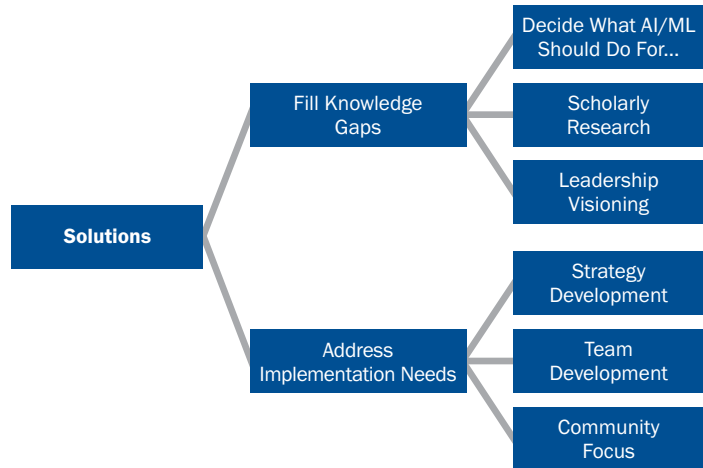
The identification of solutions with AI/ML is divided into two types of recommendations. One type is the need for the IC to be reformed: those solutions that require filling in unknowns called knowledge gaps. The other type is the need for effective action, called implementation needs. Figure 20 summarizes the two solutions and the more detailed six solutions discussed below.

Fill Knowledge Gaps

Fill Knowledge Gaps solutions consist of three types, as shown in Figure 20: making organizational decisions about what AI/ML should support (*decide what AI/ML should do for...*), identifying knowledge gaps for which scholarly research is needed to fill the gaps (*scholarly research*), and leadership-related identification of ways forward (*leadership visioning*).

Each of the three types is further divided into specific solutions, which are discussed below in more detail: *Decide What AI/ML Should Do For...* includes a focus on analysis, collection, data, a conceptual framework, low-deliberative tasks, machine-human team, priorities, and process, summarized in Figure 21; *Scholarly Research* includes AI/ML Capabilities, Fill Foundational Knowledge, Machine-Human Team, Real-World Understanding, and Theory of Use, summarized in Figure 22; and *Leadership Visioning* includes leaders develop need to change, solve a problem not buy a solution, use scientific principles, focus on adversary AI/ML, imagine new roles for AI/ML, and AI/ML augments Human, summarized in Figure 23.

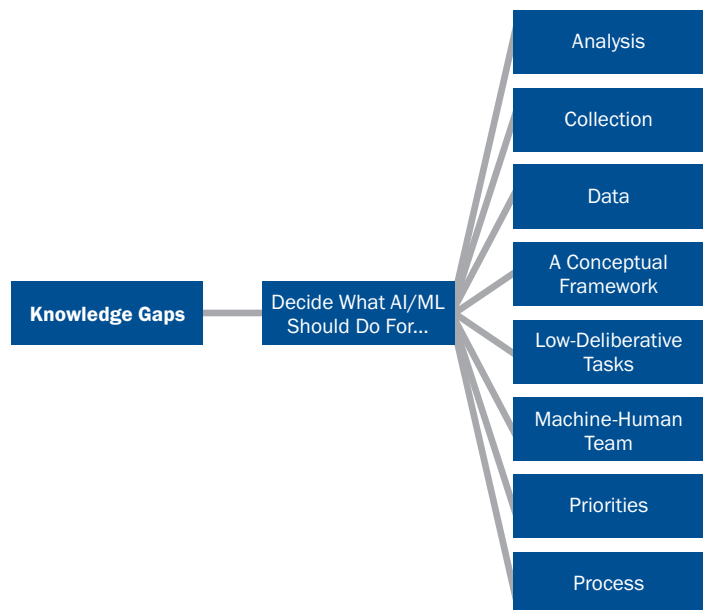
Figure 20. Solutions for AI/ML for the IC



Decide What AI/ML Should Do For...

Deciding what AI/ML should do for *analysis* includes having the capability of AI/ML algorithms to allow social media companies to remove or stop adversary's content since the tools can do it faster than analysts.²⁹⁷ Because the number of Internet-connected devices is so great and human interaction through the Internet is so vast, AI/ML algorithms should be used to maintain an advantage over analysts deluged by information overload.²⁹⁸ AI/ML algorithms can also improve information gathering and change how analysts conceptualize a threat.²⁹⁹

Figure 21. Making AI/ML Decisions Solutions



The IC needs to decide which organizations should adopt AI/ML to equalize or improve the Great Power Competition.³⁰⁰ Such allocation of AI/ML priorities will help decisionmakers use AI/ML's strengths and asymmetric advantages in pursuing a goal.³⁰¹ While AI/ML is currently used in front-end analysis of data collection, AI/ML has the potential to support more deliberative tasks for the analyst.³⁰² A counter-AI capability is needed to recognize adversarial efforts to alter or use manipulated data.³⁰³

Given domain expertise and a stable environment, AI/ML can identify anomalies that provide the intelligence analyst with a broader perspective.³⁰⁴

Deciding what AI/ML should do for *collection* includes AI/ML helping with detection and early warning by identifying imperceptible changes and detecting anomalous behavior.³⁰⁵ AI/ML can also help validate the truthfulness of data and collection sources.³⁰⁶ AI/ML could help with human intelligence operations by improving the monitoring of security and counterintelligence risks.³⁰⁷ AI/ML can help by providing faster processing of edge devices in high-risk or denied areas when connected with sensors and communication platforms.³⁰⁸ AI/ML can help improve the speed and precision of captured materials.³⁰⁹

AI/ML can help with open-source information, and within that vein, the IC should create a federated approach to applying AI/ML to open-source information.³¹⁰ Open-source intelligence should be defined as the need to address a specific intelligence requirement and process to provide insights not otherwise available from classified sources.³¹¹ IC agencies should teach new analysts that open-source information can provide novel insights.³¹² AI/ML technologies can help with planning, scheduling, and tasking collection platforms based on requirements and target type for the technical intelligence sources.³¹³

Deciding what AI/ML should do for *data* includes IC agencies creating classified training data so that algorithms have realistic training data for intelligence problems.³¹⁴ However, AI/ML should not be used if agencies do not have the needed data.³¹⁵ AI/ML can help with processing and triaging data, which otherwise takes a long time if done manually.³¹⁶ IC intelligence products should be machine-readable to be disseminated at machine speeds in machine-readable formats and to support time-sensitive tasks.³¹⁷ There are topics where AI/ML is being used against big data, including maritime security, cyber security, money laundering, multi-INT analysis, and space situational awareness.³¹⁸ Yet, IC agencies need to recognize that big data cannot fill knowledge gaps as it cannot predict causality.³¹⁹ AI/ML is suitable for data cleaning.³²⁰ Agencies, however, should broaden knowledge sources beyond big data through AI/ML applications to include study groups, think tanks, media reports, and published books.³²¹

Deciding what AI/ML should do for a *conceptual framework* includes beginning with the right metrics, which requires the organization to understand how the system will be used in a detailed way.³²² The IC needs to develop a framework consisting of several dimensions: the various AI/ML technology capabilities, the diversity of AI/ML applications in the IC, and the investment an organization should expend in time and space.³²³ IC agencies should understand the models, algorithms, and heuristics of today's AI/ML tools.³²⁴ Assessing the diversity of AI/ML applications should be categorized through four independent factors: the degree of control the agency has over the AI/ML development and its deployment; the extent of relevancy of data, power, and bandwidth assumed to be available; the knowledge of how fast AI/ML algorithms are expected to process data and provide an output; and, the degree of resilience an organization has in recovering from AI/ML failures.³²⁵ A framework should also consider how trust is established in the AI/ML system to determine the limits of the AI/ML's behavior, when it may not work, and when it will work.³²⁶

Deciding what AI/ML should do for *low-deliberative tasks* includes AI/ML can save time on low-value tasks so that the analyst can devote more time and attention to high-value activities.³²⁷ There are two ways that AI/ML

adds value: having more time at the analyst's disposal and adding more high-value tasks for the analyst to process.³²⁸ This potential advantage assumes the IC agency will decide what high-value tasks for an analyst are.³²⁹ AI/ML offers value by offloading low-value tasks meaning analysts can pursue high-value tasks, including analysis, planning, and tasks requiring creativity, communication, and collaboration.³³⁰ AI/ML is a transformative technology for low-level tasks. Still, because of the increase in the number and types of threats, AI/ML cannot continue to be a transformative technology even for low-level tasks, as it will not be able to keep pace.³³¹

Deciding what AI/ML should do for *machine-human teams* includes IC organizations having a philosophy that creates and maintains machine-human tradecraft, consisting of design policy, information technology, a flexible acquisition policy, and an agile security environment.³³² The IC must determine when AI/ML is best used to support the analyst versus when AI/ML is best used to support machine-human teaming.³³³ IC agencies should identify high-impact use cases where machine-human teaming has an important impact.³³⁴

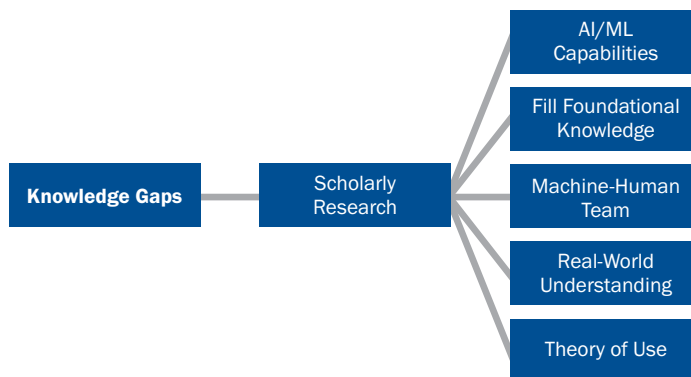
Deciding what AI/ML should do for *priorities* includes IC agencies clarifying their priorities and how AI/ML falls within that strategy.³³⁵ There should be more nuanced discussions about what AI/ML can and not do, which should include policymakers and analysts.³³⁶ The IC should prioritize where to use AI/ML at each stage of the intelligence cycle, to the extent possible, and regarding the available data.³³⁷

Deciding what AI/ML should do for *process* includes having certification labels visible on AI/ML technologies with information on its characteristics and training datasets needed.³³⁸ There is also the need for performance-tracking processes that continuously assess the AI/ML technology to evaluate its accuracy.³³⁹ To ensure quality control, there should be continuous red teaming of AI/ML models at each step.³⁴⁰ Assuming the IC has automated the AI/ML processes for each intelligence discipline, the AI/ML technology should fuse each process into an uninterrupted all-source feed.³⁴¹

Scholarly Research

To fill many of the knowledge gaps, scholarly research is needed. One area, as illustrated by Figure 22, is *AI/ML capabilities*. More research is needed to evaluate AI/ML methods so that developers and users can understand the capabilities and limitations of tools.³⁴² Research is also necessary to help understand how AI/ML can support pattern recognition, including making inferences between detected objects and visualized networks for greater clarity and understanding.³⁴³ Research is needed on having a demonstrable AI/ML tool that represents a realistic capability without simplifying the tool's ability. This would help ground future acquisition efforts and users to what can be accomplished.³⁴⁴ There

Figure 22. Scholarly Research Solutions



needs to be more research on the technical and social challenges of adopting AI/ML into the analytical work within the IC.³⁴⁵

More research is needed about how AI/ML *fills foundational knowledge*. Foundational knowledge is basic information about what exists, for example, in an adversary's military arsenal. Research into how curiosity can help fill knowledge gaps related to foundational knowledge would support generating questions that could be used to fill gaps.³⁴⁶ Investments in AI/ML include foundational basic research, so a focus on such research is warranted.³⁴⁷

Research into *machine-human teaming* covers a broad spectrum of interests. The following areas fall into the machine-human team research areas: understanding the constraints of the human analyst and the nature of how the analytic problem is represented; factors needed for trust to be built and maintained by the analysts of the machine; analyst reactions to different levels of cognition in machines; when and what the machine explains to the analyst so that the analyst can have an accurate concept of its machine team member; the ability of the machine to have perspective taking and joint attention to achieve cooperation and coordination with the analyst; understand the analyst's behavior in different machine teaming situations so that the machine has a valid model of analyst behavior; and what the machine needs to know to be able to reason with analysts.³⁴⁸ Most of these research topics would need to be accomplished through experiments.

Research into AI/ML *real-world understanding* includes communication studies in real-world contexts focused on natural language that could then be scaled up so that the machine could communicate in more complex, novel, or uncertain contexts and communication studies on the machine's ability to understand context so that the machine's communication with the analyst is grounded in the context and environment.³⁴⁹

Research into AI/ML *theory of use* involves understanding how AI/ML can be used. There needs to be more research at the micro and macro level on how the IC thinks about the entire life cycle of AI/ML, from planning to deployment to integration with other types of technologies.³⁵⁰ At the macro level, research is needed on how decisionmakers conceptualize how AI/ML could be used in their agencies.³⁵¹ Within the academic realm, the conduct of research itself must include research on cross-disciplinary collaboration.³⁵² Research is needed in AI/ML task learning for knowledge reuse and recombination across different analytic problem sets.³⁵³ There is a need for a theory of AI/ML use in the IC that includes the tasks in which AI/ML are successful and those not.³⁵⁴ Research on the future demand for AI/ML is needed to appropriately match the application to analysts who engage in complex, deliberative tasks.³⁵⁵

Leadership Visioning

The articles suggest that executive leadership must be involved in the overall challenge of incorporating AI/ML into the IC. An important step is that *leaders develop the need to change*, where leaders must create a sense of urgency because there is an AI/ML crisis, embrace out-of-the-box thinking to grow a culture of innovation, articulate acceptable risks, and expected failure rates, and back innovators when efforts fail.³⁵⁶

Creativity and commitment are needed by leadership as these are decisive qualities for achieving success in integrating and harnessing new technologies.³⁵⁷

Leadership must not succumb to the pressure to buy their way out of the AI/ML crisis, as they must *solve problems, not buy solutions*. Leadership must overcome the tendency to apply expensive solutions to a problem that may not match. Still, there is a cultural bias that the act of purchasing is an end in and of itself, a false hope strategy.³⁵⁸ Leaders need to reframe the conceptualization of AI/ML

as an information technology solution that can be bought. Instead, the technology should be conceptualized as an operational problem that needs to be solved.³⁵⁹ In support of this end, the IC should welcome the use of commercially developed AI/ML and only enter expensive modifications when necessary.³⁶⁰

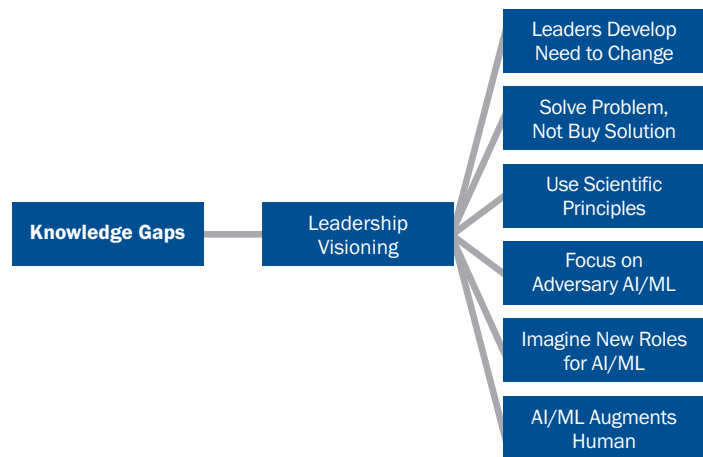
Leadership should ensure that AI/ML is grounded in *scientific principles* using data science. Data science is an interdisciplinary field integrating the domains of statistics, computing, communication, management, and sociology to study data within a context, such as particular domains but also organizational and social behaviors, to transform data into knowledge.³⁶¹ Data science must be integrated into analytic strategies against threats.³⁶² Leadership must embrace science for evidence-based decisionmaking building using the scientific method.³⁶³ Leadership must model behavior that data science is to be integrated into analysis so that strategic intelligence can be effective.³⁶⁴ In line with a foundation in science, there must be a systematic approach modeled by leadership on making risk-conscious decisions.³⁶⁵ A data science approach should frame the threat environment as a complex adaptive system.³⁶⁶

Not only must the IC understand how to use AI/ML, but it must also have a *focus on adversary AI/ML use*. Leadership must ensure the threat side is addressed. IC analysts should focus on adversarial AI/ML, which could be done by establishing cadres whose sole purpose is to focus on foreign AI/ML systems and capabilities.³⁶⁷ One way this could be accomplished is by establishing an AI/ML clearing house that would have information on foreign AI/ML use.³⁶⁸ The IC should create an AI/ML red team that would focus on the malicious use of AI/ML.³⁶⁹ By extension, such a focus would necessitate investment in counter-AI capabilities.³⁷⁰ In addition, the IC should fund adversarial AI/ML testing and red-teaming adversary AI/ML use.³⁷¹

An important leadership role is *imagining new roles for AI/ML*. AI/ML can help deliver information to decisionmakers. It can expand analysts' knowledge by bringing continuous learning to the workplace, offering AI/ML-based recommendations on what to read, and validating AI/ML tools for analysts to trust AI/ML outputs.³⁷²

Leadership must ensure the right mindset exists so that *AI/ML augments humans*, not replaces analysts. There must be agreement that intelligence is defined as primarily a human profession and that AI/ML

Figure 23. Leadership Visioning Solutions



technologies augment the profession.³⁷³ While the temptation exists to extoll AI/ML as transformative for the IC, leadership should take care to communicate AI/ML as an incremental capability and not a transformational technology.³⁷⁴ There are many cultural factors that decisionmakers should consider when seeking AI/ML technologies, such as flexibility and embracing change, a learning culture, data-driven decision-making, open communication and collaboration, shared digital vision, an entrepreneurial culture, critical thinking, and open questioning.³⁷⁵

Address Implementation Needs

Implementation needs solutions consist of three types, as shown above in Figure 24: addressing a wide variety of perspectives on how to approach facets of AI/ML (*strategy development*), addressing the unique challenges of machine-human teaming (*team development*), and leadership-related identification of ways forward (*community focus*).

Each type is further divided into specific solutions, which are discussed below in more detail. In Figure 24, *strategy development* focuses on budget, acquisition, human resources, oversight, partnership, product delivery, and risk. *Team development* includes analyst involvement, analyst understanding, interdisciplinary team, reduction of human bias, and domain expertise is summarized in Figure 25. The *community focus* includes intelligence resources, innovation, testing, and verification, as outlined in Figure 26.

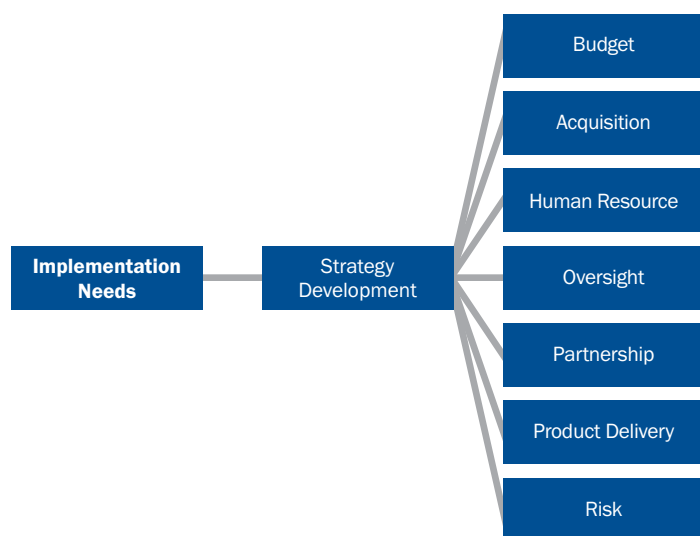
Strategy Development

The budgeting process needs repair, which drives the call for a *budget strategy*. The budgeting process must have greater IC flexibility to support AI/ML development.³⁷⁶ The budgeting process should include investments

in the technical infrastructure to support AI/ML technologies.³⁷⁷ Funding can be used as a stick but threatens to cut off funding to agencies when their AI/ML development efforts lag far behind the private sector.³⁷⁸

Similarly, the *acquisition process* needs repair. Acquisitions should use statements of objectives instead of a statement of work, as the former does not direct how the outcome can be achieved.³⁷⁹ New acquisition authorities like Other Transaction Authority (OTA) and Commercial Solutions Openings (CSO) are needed.³⁸⁰ Develop a strategy that uses a single-time budget to allow the same money to be used for various purposes.³⁸¹ Create a

Figure 24. Strategy Development Solutions



parallel organization for the acquisition of AI/ML technologies.³⁸² Similar organizations are created when the official or primary organization is not capable or designed to identify, define, or solve a strategically important problem. Yet, the parallel organization does not replace or displace the official organization.³⁸³

Developing a *human resource strategy* was a key focus area in the articles. One recommendation deals with reducing the analyst turnover that might occur due to analysts experiencing AI/ML errors or introducing new technology.³⁸⁴ Another general idea is to encourage creativity and deep thinking in analysts.³⁸⁵ Agencies need to educate IC analysts on navigating the various development paths requiring knowledge and skills in AI/ML.³⁸⁶ Agencies should incentivize being part of a project team to ensure they have a good understanding of AI/ML.³⁸⁷ Recruitment of AI/ML, including its testing and evaluation, must be accelerated, and there needs to be appropriate training once on board.³⁸⁸ The IC needs to develop trust and confidence in AI/ML by its senior leaders. Hence, the IC should develop an AI introductory course for them and identify workforce skill sets for AI-enabled tasks.³⁸⁹ There needs to be workforce skillset retooling for employees to increase the capabilities of the analytic workforce.³⁹⁰ The IC should consider rotating analysts into the AI/ML industry to learn about the AI/ML capabilities and how the industry designs and builds them.³⁹¹ Managers need to be trained in AI/ML systems before they are deployed in the workspace.³⁹²

An *oversight strategy* is needed. For one, Congressional oversight must be adaptive, where oversight is accepting of change and uncertainty.³⁹³ There must be a shared vision about AI/ML and what an AI/ML project looks like between Congress and the IC.³⁹⁴ The IC should consider adopting the DoD's DevSec-Ops oversight approach, a continuous monitoring culture, and practice to integrate software development (Dev), security (Sec), and operations (Ops).³⁹⁵ Oversight must recognize the importance of allowing for iteration, which would allow oversight to be modified if an initial approach was not working as expected.³⁹⁶ The IC and the Congressional committees must alter their interaction based on trust to manage expectations and reduce surprises.³⁹⁷ The DNI should have two informal engagements with committees, one semi-annual about AI/ML projects and the other periodically, to create mutual understanding and build trust around AI/ML projects.³⁹⁸ To build trust between the IC and committees, Congressional committees must not misuse the idea of two informal engagements.³⁹⁹ Committees should organize themselves around functions instead of agencies to achieve a deeper focus on AI/ML technologies.⁴⁰⁰

The development of an AI/ML *partnership strategy* is extensively highlighted in the articles. The IC needs to develop a plan for bridging the AI/ML innovation gap between the private sector and the IC.⁴⁰¹ AI/ML advances are being developed and experimented with in the private sector.⁴⁰² There is also the need for cooperation between the IC and law enforcement on AI/ML technologies.⁴⁰³ To have a partnership strategy, the collaboration between the private sector and IC practitioners needs to be modeled.⁴⁰⁴ Partnering with the private sector to combat online threats and take advantage of AI/ML advances is essential.⁴⁰⁵ Because of the asymmetry between the private sector and the IC, the IC should not attempt to develop its AI/ML technologies in-house.⁴⁰⁶ The IC should develop AI/ML-safety organizations that would serve as the champion for safety and be the interface with the private sector.⁴⁰⁷

Due to the potential complexities of AI/ML, a strategy is needed to think about the intelligence *product delivery* to decisionmakers. One key issue is the lack of algorithm explanation, which suggests that if the

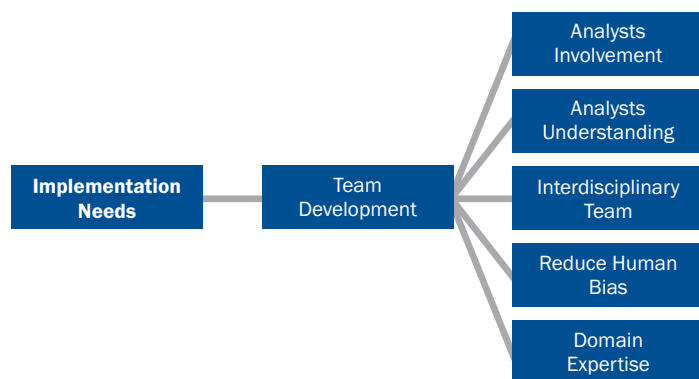
algorithm results cannot be explained, such results should be avoided.⁴⁰⁸ The IC should re-instrument the process of product delivery based on data architecture, engagement metrics, and customer modeling so that AI/ML can be applied to customer feedback.⁴⁰⁹ The IC should develop an AI/ML training course for its customers to improve trust in the IC products.⁴¹⁰ IC products should be output in human- and machine-readable formats so that results can be incorporated into other analytic efforts throughout the IC.⁴¹¹

The IC needs to develop a *risk strategy*. A flexible strategic risk framework would apply across the IC, allowing each agency to tailor the framework to its mission.⁴¹² Such an AI-based risk strategy would include how to address AI/ML failure, biased data, AI/ML adversarial attacks, supply chain problems, human mistakes, cost overruns, legal issues, and oversight issues.⁴¹³ There needs to be a balance between the risks of acting and not acting.⁴¹⁴ The IC must ensure that privacy is protected, that AI/ML does not infringe upon it, and that a risk mitigation strategy includes routine monitoring.⁴¹⁵ Understanding the risks associated with AI/ML will help the IC determine how much it can trust AI/ML results.⁴¹⁶ A key risk area is the testing and evaluation of AI/ML technologies, which behooves the IC to develop a risk-based testing and evaluation framework.⁴¹⁷

Team Development

Analyst involvement in the development process is a critical factor in the success of AI/ML tools. Agencies need to involve all players, especially the analysts, in the development and implementation stages of AI/ML planning to include being receptive to analyst observations and feedback.⁴¹⁸ Analyst involvement must go beyond their familiarization with algorithms as it must include the context in which the analyst works

Figure 25. Team Development Solutions



because organizational dynamics play into the workplace condition, thus shaping how analysts perceive the usefulness of the AI/ML tool.⁴¹⁹ It is necessary to understand the context of analysts, especially those in large bureaucratic institutions, because these contexts can shape the diffusion, adoption, acceptance, and usage of AI/ML tools.⁴²⁰ Some recommend that ICD 203 Analytic Standards⁴²¹ be used as a starting point for developing AI/ML themes of data transformation, aggregation, labeling, and display.⁴²²

Another critical factor is the *analyst's understanding* of AI/ML algorithms. Analysts should embrace AI/ML algorithms because these applications can sort data, learn from data, and respond to data.⁴²³ Analysts need to trust AI/ML output, which requires transparency and accountability for how AI/ML systems are used in practice.⁴²⁴ Knowing how algorithms work is especially important as more data is processed by AI/ML technologies.⁴²⁵ Analysts need to be transparent about the data used while training the AI/ML tool to reduce or manage bias.⁴²⁶ Reducing unintentional bias in AI/ML requires transparency in analysts seeing how the algorithms produce results.⁴²⁷ Analysts need to know how AI/ML systems are set up and used in

practice to trust them.⁴²⁸ Creating trust in the analyst means overcoming the analyst's initial doubt about an AI/ML tool.⁴²⁹

It is not just analyst involvement in AI/ML development, the *interdisciplinary team* needs to be involved. This is a joint effort by all players, including computer scientists, intelligence analysts, strategists, lawyers, and leadership visioning.⁴³⁰ Involving all players means analysts and managers so they can see and participate in the design, implementation, and adaptation of AI/ML so that the tool will more likely be accepted.⁴³¹ All players should be involved in deciding what analysts need and buy.⁴³² A framework needs to be developed, for the long term, so that experts from various fields can participate in how AI/ML deals with complex data sets and how the results of analysis of such data are implemented.⁴³³

Reducing human bias is not a simple level of effort and not a one-time effort. Having a diverse set of professionals involved in scrutinizing AI/ML brings different perspectives to various tasks such as review of processes, providing algorithmic transparency, and auditing the review of bias.⁴³⁴ The team's development of a shared understanding of each member's expertise and how each member frames an observation or problem can help identify bias.⁴³⁵ It is well-recognized that AI/ML bias is a systemic phenomenon. Hence, it requires a holistic approach to observing and reducing it.⁴³⁶

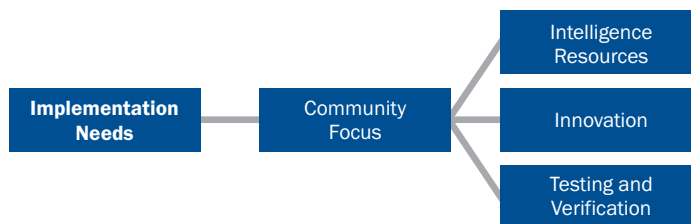
AI/ML technologies must interact with *domain expertise* because algorithms should reflect, as best as possible, the data relevant to the domain.⁴³⁷ The analyst is needed to help predefine the analytic problem space for the AI/ML algorithm. Such involvement will shape the design or selection of algorithms as it sheds light on analytic decisionmaking.⁴³⁸ Domain expertise is especially important with big data as context familiarization is needed to prioritize feasible targets.⁴³⁹

Community Focus

The solutions set of community focus involves recommendations for centralizing various functions. The first is to establish a central hub for AI/ML for *intelligence resources*. Since agencies have different missions and cultures, one important and necessary ingredient is to speak a common language about AI/ML systems, from planning to deployment.⁴⁴⁰ There are some efforts taking place, but they are not centralized. For example, there is a digital directorate at CIA, new AI/ML initiatives at NGA, and new cloud-computing efforts at NSA.⁴⁴¹ One idea is for ODNI to establish a National Artificial Intelligence Center (NAIC).⁴⁴² Another idea is to centralize AI/ML budget control with ODNI.⁴⁴³

Another idea is to establish a central hub for AI/ML *innovation*. The DNI should designate a senior innovation leader responsible for driving innovation across the IC.⁴⁴⁴ That leader could launch an IC innovation initiative.⁴⁴⁵ One idea is that a parallel organization should be established to nurture AI/ML, like Lockheed Martin's Skunk

Figure 26. Community Focus Solutions



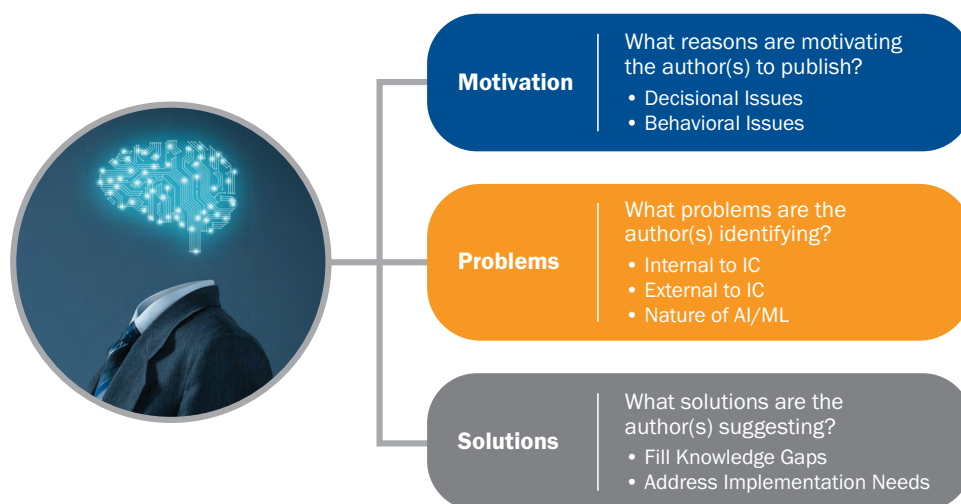
Works.⁴⁴⁶ The IC could create an unclassified sandbox for IC customers to test and evaluate new capabilities.⁴⁴⁷ Similarly, the IC could use pilot projects to prove impactful ideas demonstrating that the IC can handle flexibility and speed.⁴⁴⁸ AI/ML systems need regular maintenance, so one additional centralized function would address the fast pace of AI/ML evolution.⁴⁴⁹

There is a need for centralizing a focus on AI/ML *testing and verification*. One of the main purposes for investing in AI/ML is to strengthen testing, verification, and validation.⁴⁵⁰ A centralized body could serve as a coordinating entity to lead testing and evaluation of AI/ML technologies and to incentivize cooperation within the IC.⁴⁵¹ The centralized body could be used to translate the private sector testing framework into the context and needs of the IC.⁴⁵² This suggestion could improve trust in AI/ML but testing, training, and certifying human-machine teams through wargaming, simulation, and experimentation.⁴⁵³ The IC should create an AI/ML red cell to test and verify its use on critical threats.⁴⁵⁴

Summary of Findings

Figure 27 summarizes the answers to the three review questions posited at the beginning of this study. By and large, the reasons individuals have raised concerns about AI/ML in the IC, identified in the motivation section above, are areas that the IC does not have control over or has not done a very good job of overcoming its self-imposed constraints. The motivations fall into two broad issues: factors that should be considered in the decisionmaking process leading up to the acquisition of AI/ML and concerns about the environment within and outside of the IC about AI/ML that the IC currently has not satisfactorily mitigated the various environmental concerns. The problems identified for the IC are quantitatively many, topic-wise varied, and qualitatively very complex. Collectively, the problems identified can be overwhelming. A conceptual framework would help absorb the problems and solutions and provide a way forward. A proposed framework is discussed in the next section.

Figure 27. Answers to Review Questions



Problems and Solutions

The problems fall into three main focus areas: internal to the IC, external to the IC, and the nature of AI/ML. The problems identified as *internal to the IC* are people problems in one way or another. They primarily focus on the challenges of a culture not ready or able to incorporate new technology, innovation, and risk-taking, to name a few. However, other people-related problems occur across multiple levels of analysis, from the individual to the team, the organization, and the national security sector. Challenges range from using an outdated theory of defining intelligence to a lack of understanding that AI/ML requires an interdisciplinary approach. AI/ML is not just an information technology problem.

The problems identified *external to the IC* are varied and are also covered across multiple levels of analysis, from the organization to the national security sector, the U.S. society, and the international system. They range from actual and consequential national security threats to U.S. policies, rules, and regulations created by institutions and departments external to the IC. These problems constrain the IC's ability to quickly and effectively adapt to a world and society overwhelmed with digital technology and its data that almost everyone uses.

The *nature of AI/ML* problems includes a lack of knowledge about AI/ML and a lack of a realistic understanding of what AI/ML technologies can do. The technologies have limitations in their interactions with humans and limited ability to deal with deliberative cognitively intense tasks. AI/ML problems are primarily cognitive and span multiple levels of analysis, from the individual to the team and the organization.

The solutions offered are not one-step actions. This limit is probably the most important insight from this study. They are, by and large, a set of complex and time-consuming human endeavors that *fill knowledge gaps* requiring sophisticated change management techniques and organization. The most challenging solution is for the IC to decide what AI/ML can do across the analytic spectrum and enterprise domains. What is needed is an extensive research program into how analysts and others who use or might use AI/ML technologies think and behave in their interactions with AI/ML.

Leadership visioning will also be instrumental in the broader topic of implementation needs. Such needs would include incorporating scientific principles into the interaction relationship between analysts and AI/ML. Leadership will develop ideas about how analysts and other intelligence professionals should think about the new roles they play versus AI/ML technologies play. The more challenging part for leadership begins with strategy development. This includes the development of new policies, perceptions, and behaviors for how AI/ML technologies and analysts can work together.

Conclusion and Implications

To improve AI/ML use and integration, all the problems and solutions identified in this study are available for consideration. However, any IC choice faces considerable individual and organizational influences that can thwart the IC's ability to adapt. The study concludes with a discussion and framework of four elements that influence decisions: the multilevel nature of problems; solutions at each level; organizational attention; and ways to lead. Figure 29 is the conceptual framework for how the four elements are related.

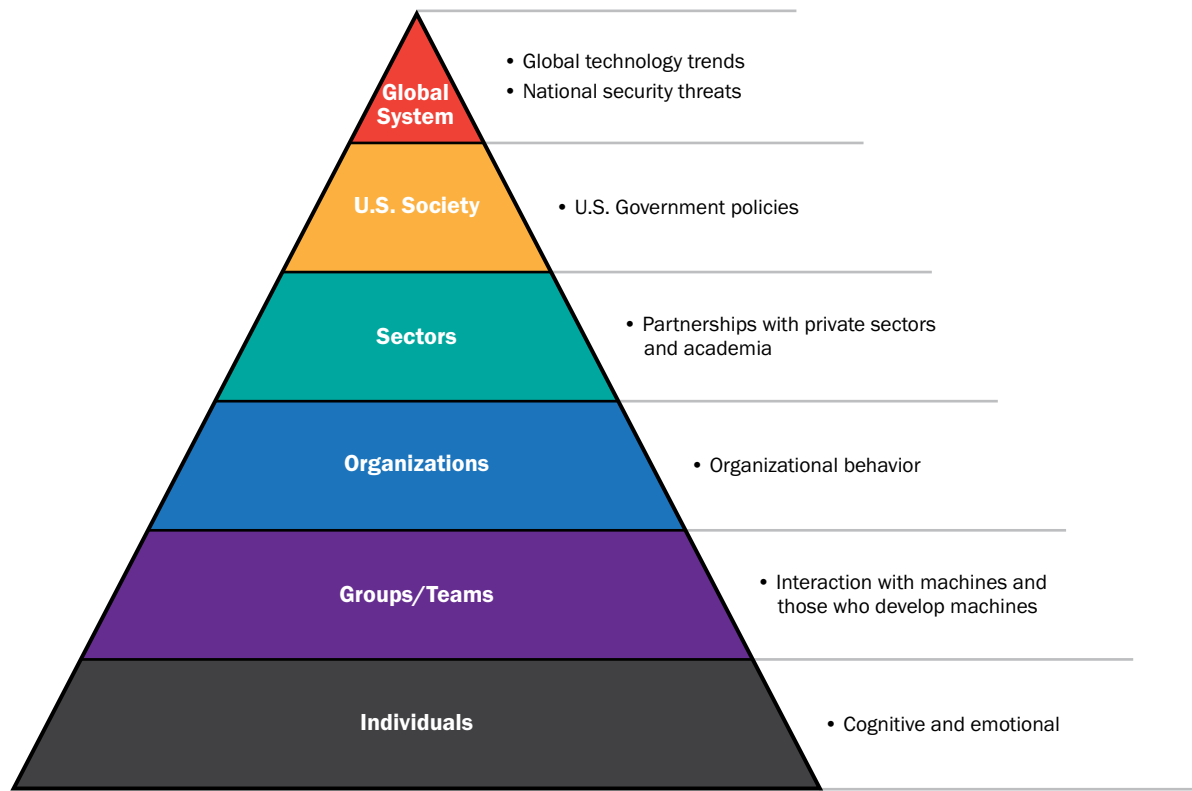
Multilevel Nature of Problems

To translate the complexity of the findings into action, a conceptual framework is informed by the theoretical lens used in this study of Ocasio's theory of organizational attention.⁴⁵⁵ The framework comprises four elements: multilevel nature of problems, solutions proposed at each level, theory of organizational attention, and ways to lead.

As to the challenges of understanding the problems identified through this study, it is helpful to represent the problems as they occur at different levels of analysis. Viewing problems from a multilevel perspective helps prioritize organizational attention on improving an understanding of the entire problem set's diversity and complexity.⁴⁵⁶

Such a multilevel perspective can be viewed as a hierarchy of different and increasingly abstract levels of the human organization represented by a triangle. Figure 28 summarizes the hierarchical view of problems from this study. The bottom of the triangle is the *individuals*. In this study, the problems at the individual level are psychological, combining cognitive and emotional factors. The next level of abstraction is behavior at the *group/team* level. Problems are mainly identified within the interaction of machines, individuals who develop the machines, and those who participate in the procurement, testing, and deployment of machines. The next level of abstraction is behavior at the organizational level. As identified in this study, problems occur throughout the IC in what is referred to as organizational behavior. The next level of abstraction is behavior at the sector level, and problems occur through partnerships with the private sector and academia. The next level is behavior at the U.S. societal level. Problems arise in the constructed policies by the U.S. Government that have a constraining force on the IC. The highest level of abstraction is behavior in the international, global system. Problems occur because of foreign actors in the international system and the proliferation of technology around the globe.

Figure 28. AI/ML in the IC as a Multilevel Human Problem Set



Solutions at Each Level

The solutions at all levels of the problem set attend to a different phenomenon.

At the global level, the IC needs a strategy to attend to the malicious use of AI/ML against the United States and its allies and to have situational awareness of global technology trends.

At the U.S. society level, strategies are needed to change how the IC is shaped by federal acquisition and budget rules, procedures, and regulations. Similarly, the IC needs to reset its relationship with the congressional oversight committees at this level.

At the sector level, scholarly research is needed from academia to address various psychological and social psychological issues within individuals and teams. There is also the need to develop a strategy to improve partnerships with academia, the private sector, and law enforcement.

At the organizational level, leadership behaviors are needed to instill a spirit for change, to change the mindset from buy-a-solution to solve-a-problem, embrace and use scientific principles, attend to the adversaries' use of AI/ML, and reinforce the proposition that AI/ML augments humans, not replace them. Strategy is also essential at this level to determine what the IC wants AI/ML to do, to make AI/ML part of the

day-to-day workflow, shift the focus of culture away from being risk-averse, and figure out how to organize everything mentioned at the organizational level of analysis in this study's findings—whether to centralize efforts or to establish a parallel organization.

At the group/team level, there is a significant learning curve needed by the various players involved in the planning, development, testing, and deployment of AI/ML, as well as the need to understand and develop mitigation processes dealing with bias.

At the individual level, many of the activities at the sector, organizational, and group/team levels need to be understood by the individual because it is the individual, whether in a team situation or not, that does the work. Almost everything at these levels needs to be grounded in understanding what the individual can and cannot do. This involves individuals acquiring knowledge and developing a strategy to prevent higher-level solutions from misplacing unrealistic expectations.

Organizational Attention

This study informs the complexity of AI/ML problems and solutions for IC decisionmakers to attend to and absorb. The various individual and organizational factors involved were discussed in the theoretical lens section of this study. Summarizing these factors, decisionmakers decide where they want to focus at their individual level. What they choose to pay attention to is shaped by the context and situation they find themselves in. The context and situation are affected by rules, resources, and social relationships. At the organizational level, attention is shaped by the organizational strategy, by the degree to which the decisionmaker takes the time and effort to engage in the organizational workflow needed to make the object of attention materialize and is affected by the history of what and how issues have been attended to in their organization.

Ways to Lead

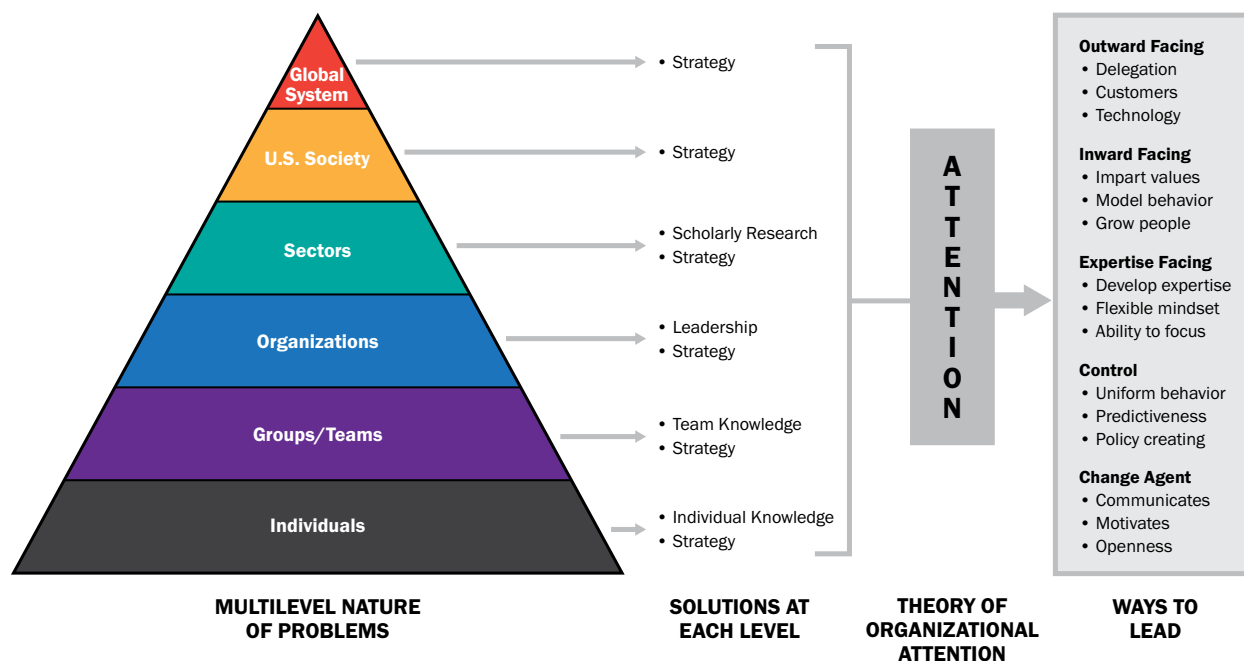
As the theory of organizational attention is centered around decisionmakers' attention, what the decisionmakers do after focusing on an issue becomes relevant for a conceptual framework. Organizational consultant Charles Farkas's empirical research into five leadership styles for implementing solutions is therefore appropriate.⁴⁵⁷ *Outward facing* is illustrated when leaders are willing to delegate much of their authority to internal organizational matters, deal with customers, and stay informed of technological advances. A second style is *inward facing*, where the leader's primary function is to impart values by modeling behavior and growing people within the organization. A third style is *expertise facing*, defined by the leader who focuses on developing expertise and looks for individuals with a flexible mindset and the ability to focus on tasks requiring expertise. A fourth style is *control*, where the leader is primarily concerned with uniform behavior in the organization, reduces uncertainty by focusing on the repeatability of processes to ensure accurate prediction of behaviors, and oversees the creation of policies to mandate control. The fifth style is the *change agent*, where the leader is primarily involved in communicating with the workforce, seeks to motivate the

workforce, and exhibits an openness to ideas for change. The styles are not mutually exclusive, and it is more likely than not that leaders will employ multiple styles across different times and contexts.⁴⁵⁸ Deciding which style to adopt for any solution, at any time, in a context, and for how long will be the key to success.

Summary of Conceptual Framework

A conceptual framework shows how the four elements are related, as shown in Figure 29. The *multilevel nature of AI/ML* problems in the IC is portrayed as an increasingly abstract set of human activities. These problems *inform proposed solutions*, each different for each level of a problem. A variety of solutions is available for a *decisionmaker's attention*. What the decisionmaker pays attention to involves not only an interest in any one solution but also a variety of factors typically not in the leader's control. These factors affect the context or situation the leader finds themselves in and the organizational policies, practices, and culture. Once attention is secured—if it is, that is—on a set of solutions, the *leader must figure out how to lead* the implementation of a solution.

Figure 29. Conceptual Framework for Addressing AI/ML in the IC



Recommendations

The conceptual framework presented in Figure 29 is not prescriptive. Instead, it is an analytic tool. Each element is multifaceted, dependent on initial starting conditions and the journey through each component in the framework.

This recommendation section is focused on a pair of conflicting solutions presented in this study, solutions that occur at a very grand scale: whether to centralize an organized effort to address all or many of the aspects of AI/ML issues in the IC, or to create a parallel organization to implement the desired new outcomes, while letting the existing system do what it currently does best without extensive disruption.

Centralization seems to be a standard approach to solving complex problems in large, hierarchical bureaucracies. On the other hand, parallel organizations are designed especially for situations where the desired outcomes are far different from the official organization's capability. The challenges of imparting a significant change effort, AI/ML, into the IC are reflected in this study's breadth and depth of problems and solutions. Almost all the problems and solutions are about people, even though the topic is technology. As such, it would seem difficult for existing organizations to solve novel and possibly intractable organizational problems.

The creation of a parallel organization would not be a centralized entity but rather a separate but connected organizational unit that operates much differently than the official system. Well-known examples of parallel organizations include Lockheed's Skunk Works, SWAT teams within police departments, and special operations forces within the U.S. Army. The set of President's Daily Brief (PDB) briefers is a parallel organization.⁴⁵⁹ Such an AI/ML-focused IC entity would likely have a better chance of overcoming the problems identified in this study if the origins of the problems were mitigated or eliminated and the solutions were supported. While a parallel organization would have cultural challenges, the opportunities for success would outweigh such challenges as seen from the success of the parallel organizations mentioned above. For example, there may be resentment in the official system for those operating in the parallel organization. Yet, the satisfaction of customers and decisionmakers would measure the parallel organization's critical success factors.

Limitations

This project has two main limitations with using the systematic review methodology. First, the three-month time frame meant that not every step in the methodology could be accomplished as rigorously as desired. Additional time would have likely increased the number of articles considered for addressing AI/ML use in the IC. Second, because of the short time frame, a limited number of databases were searched, and a limited number of search phrases were used. More databases or knowledge stores could be explored with additional time, and further consideration of search phrases may have increased articles of interest.

Concluding Comments

This systematic review of the literature about AI/ML in the IC produced a large diversity of views on problems and solutions. It is much easier to identify problems in the present than to predict solutions in the future, hence the biggest value of this study is a deeper understanding of the current state of AI/ML in the IC. This study led to four surprises. First, if decisionmakers only focus on topics of interest to them,

they would not see the entire picture or dependencies between each problem and solution. On the other hand, knowing that the complexity of the problems and solutions can overwhelm decisionmakers, the conceptual framework is constructed to alleviate this concern as much as possible. The second surprise is that almost all of the articles address both AI/ML and the IC present problems and solutions primarily in terms of human behavior, not technological or engineering issues. The third surprise is that problems and solutions about AI/ML in the IC occur at every level of analysis of the human enterprise, from the individual to the international system. Finally, the systematic review methodology provides a useful approach to intelligence studies. While the methodology has a successful history of supporting the medical, public health, environment, and educational fields, it has only recently contributed to the for-profit management field. There are few examples of using a systematic review for national security contexts, particularly the Intelligence Community.

Endnotes

1. For example, Cortney Weinbaum, Bradley Knopp, Soo Kim, and Yuliya Shokh, “Options for Strengthening All-Source Intelligence.” (Santa Monica, CA: RAND Corporation, 2022). <https://apps.dtic.mil/sti/pdfs/AD1161438.pdf>
2. Kwasi Mitchell et al., “The Future of Intelligence Analysis: A Task-Level View of the Impact of Artificial Intelligence on Intelligence Analysis,” *Deloitte Insights*, (Washington, DC: Deloitte, 2019), https://www2.deloitte.com/content/dam/insights/us/articles/6306_future-of-intel-analysis/DI_Future-of-intel-analysis.pdf
3. Avril Haines and Stephanie O’Sullivan, “Maintaining the Intelligence Edge: Reimagining and Reinventing Intelligence through Innovation.” (Washington, DC: Center for Strategic and International Studies, January 2021). https://csis-website-prod.s3.amazonaws.com/s3fs-public/publication/210113_Intelligence_Edge.pdf
4. Danielle Tarraf et al., “The Department of Defense Posture for Artificial Intelligence: Assessment and Recommendations,” (Santa Monica, CA: RAND Corporation, 2019). https://www.rand.org/pubs/research_reports/RR4229.html
5. Eric Schmidt and Bob Work, “Final Report: National Security Commission on Artificial Intelligence (AI),” (Washington, DC: National Security Commission on Artificial Intelligence, 2021). <https://apps.dtic.mil/sti/pdfs/AD1124333.pdf>
6. Michele Flournoy, Avril Haines, and Gabrielle Chetifz, “Building Trust through Testing: Adapting DoD’s Test and Evaluation, Validation, and Verification (TEVV) Enterprise for Machine Learning Systems, including Deep Learning Systems,” (Washington, DC: Georgetown University Press, 2020). <https://cset.georgetown.edu/wp-content/uploads/Building-Trust-Through-Testing.pdf>
7. Karen Jani, “The Promise and Prejudice of Big Data in Intelligence Community,” (Atlanta, GA: Georgia Institute of Technology Press, 2016). <https://arxiv.org/pdf/1610.08629v1.pdf>
8. William Ocasio, “Towards an Attention-Based View of the Firm.” *Strategic Management Review* 18 (1997): 187-206. [https://doi.org/10.1002/\(SICI\)1097-0266\(199707\)18:1+%3C187::AID-SMJ936%3E3.0.CO;2-K](https://doi.org/10.1002/(SICI)1097-0266(199707)18:1+%3C187::AID-SMJ936%3E3.0.CO;2-K)
9. William Ocasio, “Attention to Attention.” *Organization Science* 22, no. 5 (2011): 1286-1296. <https://doi.org/10.1287/orsc.1100.0602>
10. Celine Rojon, Adun Okupe, and Almuth McDowall, “Utilization and Development of Systematic Reviews in Management Research: What Do We Know and Where Do We Go from Here?” *International Journal of Management Reviews* 23, no. 2 (2021): 191-223. <https://doi.org/10.1111/ijmr.12245>
11. Yu Xiao and Maria Watson, “Guidance on Conducting a Systematic Literature Review,” *Journal of Planning Education and Research* 39, no. 1 (2019): 93-112. <https://doi.org/10.1177%2F0739456X17723971>
12. Valerie Smith et al., “Methodology in Conducting a Systematic Review of Systematic Reviews of Healthcare Interventions.” *BMC Medical Research Methodology* 11 (2011): 1-6. <https://bmcmedresmethodol.biomedcentral.com/track/pdf/10.1186/1471-2288-11-15.pdf>
13. Xiao and Watson, “Guidance on Conducting a Systematic Literature Review.”
14. Denise Rousseau, Joshua Manning, and David Denyer, “Evidence in Management and Organizational Science: Assembling the Field’s Full Weight of Scientific Knowledge Through Synthesis,” *Academy of Management Annals* 2, no. 1 (2008): 475-515. <https://doi.org/10.5465/19416520802211651>
15. Johnny Saldaña 2021, *The Coding Manual for Qualitative Researchers*, 4th edition (Thousand Oaks, CA: SAGE, 2021).

16. Saldaña, *The Coding Manual for Qualitative Researchers*.
17. Saldaña, *The Coding Manual for Qualitative Researchers*.
18. Jay T. Knippen and Thad B. Green, "Problem Solving," *Journal of Workplace Learning* 9, no. 3 (1997): 98. <http://dx.doi.org/10.1108/13665629710164904>
19. Rob Briner, David Denyer, and Denise M. Rousseau, "Evidence Based Management: Concept Cleanup Time?" *Academy of Management Perspectives* 24, no.4 (2009): 19-32. <https://doi.org/10.5465/amp.23.4.19>
20. Saldaña, *The Coding Manual for Qualitative Researchers*.
21. Guadys L. Sanclemente, "Reliability: Understanding Cognitive Human Bias in Artificial Intelligence for National Security and Intelligence Analysis," *Security Journal* 1-21. <https://doi.org/10.1057/s41284-021-00324-z>
22. Sanclemente, "Reliability."
23. Boaz Ganor, "Artificial or Human: A New Era of Counterterrorism Intelligence." *Studies in Conflict & Terrorism* 44, no. 7 (2021): 605-624. <https://doi.org/10.1080/1057610X.2019.1568815>
24. Danielle Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
25. Daniel Ish, Jared Ettinger, and Christopher Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis," (Santa Monica: RAND Corporation, 2021). <https://apps.dtic.mil/sti/pdfs/AD1146077.pdf>
26. Micah Musser and Ashton Garriott, "Machine Learning and Cybersecurity: Hype and Reality" (Washington, DC: Georgetown University Center for Security and Emerging Technology, 2021). <https://cset.georgetown.edu/wp-content/uploads/Machine-Learning-and-Cybersecurity.pdf>
27. Lilian Alessa et al., "Surprise and Suspense: How the Intelligence Community Forgot the Future," *The International Journal of Intelligence, Security, and Public Affairs* 23, no. 3 (2021): 310-342. <https://doi.org/10.1080/23800992.2021.2006954>
28. Rocco J. Blais and Adam M. Jungdahl, "Artificial Intelligence in a Human Intelligence World," *American Intelligence Journal* 36, no. 1 (2019): 108-113. <https://www.jstor.org/stable/27066342>
29. Alex Wilner, Casey Babb, and Jessica Davis, "Four Things to Consider on the Future of AI-Enabled Deterrence," *Lawfare* (July 25, 2021). <https://www.lawfareblog.com/four-things-consider-future-ai-enabled-deterrence#>
30. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
31. Erik Blasch et al., "Artificial Intelligence Strategies for National Security and Safety Standards," Paper presented at the AAAI FSS-19, Artificial Intelligence in Government and Public Sector conference, Arlington, VA, November 7-9, 2019. <https://arxiv.org/abs/1911.05727>.
32. Alessa et al., "Surprise and Suspense."
33. Patrick Bury and Michael Chertoff, "New Intelligence Strategies for a New Decade," *The RUSI Journal* 165, no. 4 (2020): 42-53. <https://doi.org/10.1080/02684527.2018.1452555>
34. David R. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition." *The Heritage Foundation*, November 17, 2020. <https://www.heritage.org/military-strength-topical-essays/2021-essays/the-intelligence-posture-america-needs-age-great-power>
35. David H. McCormick, Charles E. Luftig, and James Cunningham. "Economic Might, National Security, and the Future of American Statecraft." *Texas National Security Review* 3, no. 3 (2020): 50-75. https://repositories.lib.utexas.edu/bitstream/handle/2152/83223/04_TNSRVol3Issue3McCormick.pdf?sequence=2
36. Mitchell et al., "The Future of Intelligence Analysis."
37. Aaron F. Brantly, "When Everything Becomes Intelligence: Machine Learning and the Connected World." *Intelligence and National Security* 33, no. 4 (2018): 562-573. <https://doi.org/10.1080/02684527.2018.1452555>
38. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
39. Corin R. Stone, "Artificial Intelligence in the Intelligence Community: Culture is Critical." *Just Security*, August 17, 2021. <https://www.justsecurity.org/77783/artificial-intelligence-in-the-intelligence-community-culture-is-critical/>

40. Coring R. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition." *Just Security*, September 28, 2021. <https://www.justsecurity.org/78362/artificial-intelligence-in-the-intelligence-community-the-tangled-web-of-budget-acquisition/>
41. Corin R. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward." *Just Security*, October 19, 2021. <https://www.justsecurity.org/78641/artificial-intelligence-in-the-intelligence-community-know-risk-know-reward/>
42. Corin R. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough." *Just Security*, July 12, 2021. <https://www.justsecurity.org/77354/artificial-intelligence-in-the-intelligence-community-money-is-not-enough/>
43. Sunday O. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies." *Journal of Strategic Security* 12, no. 1 (2019): 72-106. <https://www.jstor.org/stable/26623078?seq=1>
44. Cortney Weinbaum and John N.T. Shanahan, "Intelligence in a Data-Driven Age." *Joint Forces Quarterly* 90, no. 3 (2018): 4-9. https://ndupress.ndu.edu/Portals/68/Documents/jfq/jfq-90/jfq-90_4-9_Weinbaum-Shanahan.pdf
45. Brian Katz, "The Collection Edge: Harnessing Emerging Technologies for Intelligence Collection." *Center for Strategic and International Studies*. July 1, 2020. <https://www.jstor.org/stable/resrep25236?seq=1>
46. Blais and Jungdahl, "Artificial Intelligence in a Human Intelligence World."
47. Mitchell et al., "The Future of Intelligence Analysis."
48. Greg Allen and Taneiel Chan, "Artificial Intelligence and National Security." Harvard University, July 2017. <https://www.belfercenter.org/sites/default/files/files/publication/AI%20NatSec%20-%20final.pdf>
49. Blais and Jungdahl, "Artificial Intelligence in a Human Intelligence World."
50. McCormick, Luftig, and Cunningham. "Economic Might, National Security, and the Future of American Statecraft."
51. Amy Zegart, "9/11: Look Back and Learn." *Hoover Digest*, January 20, 2020. <https://www.hoover.org/research/911-look-back-and-learn>
52. Amy Zegart, "American Spy Agencies Are Struggling in the Age of Data." *Wired*, February 2, 2022. <https://www.wired.com/story/spies-algorithms-artificial-intelligence-cybersecurity-data/#:~:text=In%20short%2C%20data%20volume%20and,struggling%20to%20adapt%20to%20it.>
53. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
54. Schmidt and Work, "Final Report."
55. Zegart, "American Spy Agencies Are Struggling in the Age of Data."
56. Emily Harding, "Move Over JARVIS, Meet OSCAR." *Center for Strategic and Intelligence Studies*, January 2022. https://csis-website-prod.s3.amazonaws.com/s3fs-public/publication/220119_Harding_MoveOverJARVIS_MeetOSCAR_0.pdf?NqfrbU05ULzzcySzNHB0pTzsNYw3HdfK
57. Laura McNamara, "Innovation Adoption, Technology Acceptance, and Data Science: Why Algorithmic Technologies are so Tricky." Slides present at the Department of Defense Human Systems Integration Community of Practice Seminar Series, Albuquerque, NM, September 4, 2020.
58. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
59. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
60. Douglas Yeung, "When AI Misjudgment Is Not an Accident." *Scientific American*, October 19, 2018. <https://blogs.scientificamerican.com/observations/when-ai-misjudgment-is-not-an-accident/>
61. Rena DeHenre, "Why the Department of Defense Should Create an AI Red Team." *Over the Horizon*, September 7, 2021. <https://othjournal.com/2021/09/07/why-the-department-of-defense-should-create-an-ai-red-team/>

62. Patrick Tucker, "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community." *Defense One*, January 27, 2020. <https://www.defenseone.com/technology/2020/01/spies-ai-future-artificial-intelligence-us-intelligence-community/162673/>
63. Zigmund Hampel-Arias, and John S. Meyers, "What AI Can and Cannot Do for the Intelligence Community." *Defense One*, January 5, 2021. <https://www.defenseone.com/ideas/2021/01/what-ai-can-and-cannot-do-intelligence-community/171195/>
64. Douglas Yeung et al., "Identifying Systemic Bias in the Acquisition of Machine Learning Decision Aids for Law Enforcement Applications," 2021, RAND Corporation, <https://www.rand.org/pubs/perspectives/PEA862-1.html>
65. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
66. Peter J. Phillips, and Gabriela Pohl. "Countering Intelligence Algorithms: Decision Theory, Design Choices & Counter-AI." *The RUSI Journal* 165, no. 7 (2020): 22-32. <https://doi.org/10.1080/03071847.2021.1893126>
67. Brantly, "When Everything Becomes Intelligence."
68. Kathleen M. Vogel et al., "The Impact of AI on Intelligence Analysis: Tackling Issues of Collaboration, Algorithmic Transparency, Accountability, and Management." *Intelligence and National Security* 36, no. 6 (2021): 827-848. <https://doi.org/10.1080/02684527.2021.1946952>
69. Kathleen M. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce: A Research Agenda." *IEEE Technology and Society Magazine* 40, no. 3 (2021): 84-92. <https://doi.org/10.1109/MTS.2021.3104384>
70. Laura A. McNamara, "Interdisciplinary Research in the National Laboratories." In *Anthropologists in the Security-Scape*, edited by Robert Albrow, George March, Laura A. McNamara, and Monica Schoch-Spana, 87-100. (Walnut Creek, NY: Left Coast Press, 2012).
71. John Laird, Charan Ranganath, and Samuel Gershman, "Future Directions in Human Machine Teaming Workshop." Report prepared at the Future Directions in Human Machine Teaming Workshop series, Arlington, VA, July 16-17, 2019. <https://basicresearch.defense.gov/Portals/61/Future%20Directions%20in%20Human%20Machine%20Teaming%20Workshop%20report%20%20%28for%20public%20release%29.pdf>
72. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
73. Schmidt and Work, "Final Report."
74. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
75. Wilner, Babb, and Davis, "Four Things to Consider on the Future of AI-Enabled Deterrence."
76. Corin R. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough." *Just Security*, July 12, 2021. <https://www.justsecurity.org/77354/artificial-intelligence-in-the-intelligence-community-money-is-not-enough/>
77. Zegart, "9/11: Look Back and Learn."
78. Zegart, "9/11: Look Back and Learn."
79. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
80. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
81. Stone, "Artificial Intelligence in the Intelligence Community: Culture is Critical."
82. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward."
83. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough."
84. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward."
85. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
86. Blais and Jungdahl, "Artificial Intelligence in a Human Intelligence World."
87. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
88. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward."
89. Harding, "Move Over JARVIS, Meet OSCAR."
90. Harding, "Move Over JARVIS, Meet OSCAR."

91. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough."
92. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward."
93. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
94. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough."
95. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
96. Alessa et al., "Surprise and Suspense."
97. Brantly, "When Everything Becomes Intelligence."
98. Alessa et al., "Surprise and Suspense."
99. Brantly, "When Everything Becomes Intelligence."
100. Brantly, "When Everything Becomes Intelligence."
101. Alessa et al., "Surprise and Suspense."
102. Alessa et al., "Surprise and Suspense."
103. Brantly, "When Everything Becomes Intelligence."
104. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
105. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
106. Alessa et al., "Surprise and Suspense."
107. Alessa et al., "Surprise and Suspense."
108. Alessa et al., "Surprise and Suspense."
109. Alessa et al., "Surprise and Suspense."
110. Brantly, "When Everything Becomes Intelligence."
111. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
112. Alessa et al., "Surprise and Suspense."
113. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
114. Alessa et al., "Surprise and Suspense."
115. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
116. Vogel et al., "The Impact of AI on Intelligence Analysis."
117. Sanclemente, "Reliability."
118. Vogel et al., "The Impact of AI on Intelligence Analysis."
119. Harding, "Move Over JARVIS, Meet OSCAR."
120. Harding, "Move Over JARVIS, Meet OSCAR."
121. Harding, "Move Over JARVIS, Meet OSCAR."
122. Laura A. McNamara, "Interdisciplinary Research in the National Laboratories."
123. Laura A. McNamara, "Interdisciplinary Research in the National Laboratories."
124. Vogel, Reid, Kampe, and Jones. "The Impact of AI on Intelligence Analysis."
125. Mitchell et al., "The Future of Intelligence Analysis."
126. Harding, "Move Over JARVIS, Meet OSCAR."
127. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce."
128. Blasch et al., "Artificial Intelligence Strategies for National Security and Safety Standards."
129. Vogel et al., "The Impact of AI on Intelligence Analysis."
130. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
131. Ganor, "Artificial or Human."
132. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce."

133. Laura A. McNamara, "Interdisciplinary Research in the National Laboratories."
134. Harding, "Move Over JARVIS, Meet OSCAR."
135. Allen and Chan, "Artificial Intelligence and National Security."
136. Zegart, "American Spy Agencies Are Struggling in the Age of Data."
137. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
138. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
139. Harding, "Move Over JARVIS, Meet OSCAR."
140. Katz, "The Collection Edge."
141. Sanclemente, "Reliability."
142. Phillips and Pohl, "Countering Intelligence Algorithms."
143. Blais and Jungdahl, "Artificial Intelligence in a Human Intelligence World."
144. Vogel et al., "The Impact of AI on Intelligence Analysis."
145. Zegart, "9/11: Look Back and Learn."
146. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
147. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
148. Musser and Garriott, "Machine Learning and Cybersecurity."
149. Harding, "Move Over JARVIS, Meet OSCAR."
150. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
151. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
152. Mitchell et al., "The Future of Intelligence Analysis."
153. Harding, "Move Over JARVIS, Meet OSCAR."
154. McNamara, "Innovation Adoption, Technology Acceptance, and Data Science."
155. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
156. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
157. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
158. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
159. Blasch et al., "Artificial Intelligence Strategies for National Security and Safety Standards."
160. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
161. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
162. Katz, "The Collection Edge."
163. McNamara, "Innovation Adoption, Technology Acceptance, and Data Science."
164. Don Fallis, "The Varieties of Disinformation." In *The Philosophy of Information Quality*, edited by Luciano Floridi and Phyllis Illari, 135-162. (New York, NY: Springer, 2014).
165. Andrew Chadwick and James Stanyer, "Deception as a Bridging Concept in the Study of Disinformation, Misinformation, and Misperceptions: Toward a Holistic Framework." *Communication Theory*, 32, no. 1 (2014): 1-24. <https://doi.org/10.1093/ct/qtab019>
166. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
167. Zegart, "9/11: Look Back and Learn."
168. DeHenre, "Why the Department of Defense Should Create an AI Red Team."
169. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
170. Katz, "The Collection Edge."
171. Schmidt and Work, "Final Report."
172. Katz, "The Collection Edge."

173. Allen and Chan, "Artificial Intelligence and National Security."
174. Schmidt and Work, "Final Report."
175. Yeung, "When AI Misjudgment Is Not an Accident."
176. DeHenre, "Why the Department of Defense Should Create an AI Red Team."
177. Alessa et al., "Surprise and Suspense."
178. Schmidt and Work, "Final Report."
179. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
180. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
181. Stone, "Artificial Intelligence in the Intelligence Community: Culture is Critical."
182. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
183. Zegart, "9/11: Look Back and Learn."
184. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
185. Zegart, "9/11: Look Back and Learn."
186. Wilner, Babb, and Davis, "Four Things to Consider on the Future of AI-Enabled Deterrence."
187. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce."
188. Harding, "Move Over JARVIS, Meet OSCAR."
189. Schmidt and Work, "Final Report."
190. Zegart, "American Spy Agencies Are Struggling in the Age of Data."
191. Katz, "The Collection Edge."
192. McCormick, Luftig, and Cunningham, "Economic Might, National Security, and the Future of American Statecraft."
193. DeHenre, "Why the Department of Defense Should Create an AI Red Team."
194. Blais and Jungdahl, "Artificial Intelligence in a Human Intelligence World."
195. Harding, "Move Over JARVIS, Meet OSCAR."
196. Corin R. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight." *Just Security*, November 30, 2021. <https://www.justsecurity.org/79254/artificial-intelligence-in-the-intelligence-community-oversight-must-not-be-an-oversight/>
197. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough."
198. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
199. Harding, "Move Over JARVIS, Meet OSCAR."
200. DeHenre, "Why the Department of Defense Should Create an AI Red Team."
201. Katz, "The Collection Edge."
202. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition."
203. Harding, "Move Over JARVIS, Meet OSCAR."
204. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
205. Brantly, "When Everything Becomes Intelligence."
206. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
207. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
208. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
209. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
210. Zegart, "American Spy Agencies Are Struggling in the Age of Data."

211. Brantly, "When Everything Becomes Intelligence."
212. Sanclemente, "Reliability."
213. Tucker, "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community."
214. Zegart, "American Spy Agencies Are Struggling in the Age of Data."
215. Schmidt and Work, "Final Report."
216. Mitchell et al., "The Future of Intelligence Analysis."
217. Harding, "Move Over JARVIS, Meet OSCAR."
218. Mitchell et al., "The Future of Intelligence Analysis."
219. Mitchell et al., "The Future of Intelligence Analysis."
220. Katz, "The Collection Edge."
221. Musser and Garriott, "Machine Learning and Cybersecurity."
222. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
223. Zegart, "American Spy Agencies Are Struggling in the Age of Data."
224. Wilner, Babb, and Davis, "Four Things to Consider on the Future of AI-Enabled Deterrence."
225. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
226. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
227. DeHenre, "Why the Department of Defense Should Create an AI Red Team."
228. Katz, "The Collection Edge."
229. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
230. Mitchell, Mariani, Routh, Keyal, and Mirkow, "The Future of Intelligence Analysis."
231. Tucker, "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community."
232. Ganor, "Artificial or Human."
233. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
234. McNamara, "Innovation Adoption, Technology Acceptance, and Data Science."
235. Ganor, "Artificial or Human."
236. Mitchell et al., "The Future of Intelligence Analysis."
237. Vogel, Reid, Kampe, and Jones. "The Impact of AI on Intelligence Analysis."
238. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
239. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
240. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
241. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
242. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
243. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
244. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
245. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
246. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
247. Brantly, "When Everything Becomes Intelligence."
248. Blaschet et al., "Artificial Intelligence Strategies for National Security and Safety Standards."
249. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
250. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
251. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
252. Katz, "The Collection Edge."
253. Tucker, "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community."

- 254. Ish, Ettinger, and Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis."
- 255. Ish, Ettinger, and Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis."
- 256. Ish, Ettinger, and Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis."
- 257. Katz, "The Collection Edge."
- 258. Tucker, "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community."
- 259. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
- 260. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
- 261. Tucker, "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community."
- 262. Brantly, "When Everything Becomes Intelligence."
- 263. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
- 264. Sanclemente, "Reliability."
- 265. Yeung, "When AI Misjudgment Is Not an Accident."
- 266. Sanclemente, "Reliability."
- 267. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
- 268. Sanclemente, "Reliability."
- 269. Phillips and Pohl, "Countering Intelligence Algorithms."
- 270. Phillips and Pohl, "Countering Intelligence Algorithms."
- 271. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
- 272. Sanclemente, "Reliability."
- 273. Vogel et al., "The Impact of AI on Intelligence Analysis."
- 274. McNamara, "Innovation Adoption, Technology Acceptance, and Data Science."
- 275. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
- 276. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
- 277. Phillips and Pohl, "Countering Intelligence Algorithms."
- 278. Phillips and Pohl, "Countering Intelligence Algorithms."
- 279. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
- 280. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
- 281. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
- 282. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
- 283. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
- 284. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
- 285. Blasch et al., "Artificial Intelligence Strategies for National Security and Safety Standards."
- 286. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
- 287. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
- 288. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
- 289. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
- 290. Phillips and Pohl, "Countering Intelligence Algorithms."
- 291. Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
- 292. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
- 293. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."

294. Mitchell et al., "The Future of Intelligence Analysis."
295. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
296. Wilner, Babb, and Davis, "Four Things to Consider on the Future of AI-Enabled Deterrence."
297. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
298. Brantly, "When Everything Becomes Intelligence."
299. McCormick, Luftig, and Cunningham, "Economic Might, National Security, and the Future of American Statecraft."
300. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
301. McCormick, Luftig, and Cunningham, "Economic Might, National Security, and the Future of American Statecraft."
302. Katz, "The Collection Edge."
303. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
304. Wilner, Babb, and Davis, "Four Things to Consider on the Future of AI-Enabled Deterrence."
305. Katz, "The Collection Edge."
306. Katz, "The Collection Edge."
307. Katz, "The Collection Edge."
308. Katz, "The Collection Edge."
309. Katz, "The Collection Edge."
310. Schmidt and Work, "Final Report."
311. Harding, "Move Over JARVIS, Meet OSCAR."
312. Harding, "Move Over JARVIS, Meet OSCAR."
313. Katz, "The Collection Edge."
314. Tucker, "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community."
315. Mitchell et al., "The Future of Intelligence Analysis."
316. Katz, "The Collection Edge."
317. Schmidt and Work, "Final Report."
318. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
319. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
320. Mitchell et al., "The Future of Intelligence Analysis."
321. Alessa et al., "Surprise and Suspense."
322. Ish, Ettinger, and Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis."
323. Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
324. Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
325. Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
326. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
327. Mitchell et al., "The Future of Intelligence Analysis."
328. Mitchell et al., "The Future of Intelligence Analysis."
329. Mitchell et al., "The Future of Intelligence Analysis."
330. Mitchell et al., "The Future of Intelligence Analysis."
331. Musser and Garriott, "Machine Learning and Cybersecurity."
332. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
333. Schmidt and Work, "Final Report."
334. Laird, Ranganath, and Gershman. "Future Directions in Human Machine Teaming Workshop."
335. Mitchell et al., "The Future of Intelligence Analysis."
336. Musser and Garriott, "Machine Learning and Cybersecurity."

337. Schmidt and Work, "Final Report."
338. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
339. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
340. Yeung et al., "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications."
341. Schmidt and Work, "Final Report."
342. Ish, Ettinger, and Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis."
343. Katz, "The Collection Edge."
344. Wilner, Babb, and Davis, "Four Things to Consider on the Future of AI-Enabled Deterrence."
345. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce."
346. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
347. Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
348. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
349. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
350. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce."
351. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce."
352. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
353. Laird, Ranganath, and Gershman, "Future Directions in Human Machine Teaming Workshop."
354. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
355. Hampel-Arias, and Meyers, "What AI Can and Cannot Do for the Intelligence Community."
356. Stone, "Artificial Intelligence in the Intelligence Community: Culture is Critical."
357. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
358. Alessa, Moon, Valentine, Works, Hepburn, and Kilskey, "Surprise and Suspense."
359. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
360. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
361. Longbing Cao, "Data Science: A Comprehensive Overview." *ACM Computing Surveys*, 50, no. 3 (2017): 1-42. <https://doi.org/10.1145/3076253>
362. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
363. Alessa et al., "Surprise and Suspense."
364. Bury and Chertoff, "New Intelligence Strategies for a New Decade."
365. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward."
366. Alessa, Moon, Valentine, Works, Hepburn, and Kilskey, "Surprise and Suspense."
367. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
368. Blais and Jungdahl, "Artificial Intelligence in a Human Intelligence World."
369. DeHenre, "Why the Department of Defense Should Create an AI Red Team."
370. Allen and Chan, "Artificial Intelligence and National Security."
371. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
372. Mitchell, Mariani, Routh, Keyal, and Mirkow, "The Future of Intelligence Analysis."
373. Alessa, Moon, Valentine, Works, Hepburn, and Kilskey, "Surprise and Suspense."
374. Musser and Garriott, "Machine Learning and Cybersecurity."
375. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."

376. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough."
377. Danielle Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
378. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
379. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition."
380. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition."
381. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition."
382. Harding, "Move Over JARVIS, Meet OSCAR."
383. Thomas F. Hawk and Dale E. Zand, "Parallel Organization: Policy Formulation, Learning, and Interdivision Integration." *Journal of Applied Behavioral Science*, 50, no. 3 (2014): 307-336. <https://doi.org/10.1177%2F0021886313509276>
384. Sanclemente, "Reliability."
385. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
386. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition."
387. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition."
388. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
389. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
390. Schmidt and Work, "Final Report."
391. Harding, "Move Over JARVIS, Meet OSCAR."
392. Vogel et al., "The Impact of AI on Intelligence Analysis."
393. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough."
394. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
395. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
396. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
397. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
398. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
399. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
400. Stone, "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight."
401. Shedd, "The Intelligence Posture America Needs in an Age of Great-Power Competition."
402. Stone, "Artificial Intelligence in the Intelligence Community: Money is Not Enough."
403. Ogunlana, "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies."
404. Stone, "Artificial Intelligence in the Intelligence Community: Culture is Critical."
405. Zegart, "American Spy Agencies Are Struggling in the Age of Data."
406. Zegart, "American Spy Agencies Are Struggling in the Age of Data."
407. Allen and Chan, "Artificial Intelligence and National Security."
408. Ganor, "Artificial or Human."
409. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
410. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
411. Schmidt and Work, "Final Report."
412. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward."
413. Stone, "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward."
414. Schmidt and Work, "Final Report."
415. Ganor, "Artificial or Human."
416. Sanclemente, "Reliability."

417. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
418. Vogel, Reid, Kampe, and Jones, "The Impact of AI on Intelligence Analysis."
419. McNamara, "Interdisciplinary Research in the National Laboratories."
420. McNamara, "Innovation Adoption, Technology Acceptance, and Data Science."
421. Office of the Director of National Intelligence, *Intelligence Community Directive 203, Analytic Standards, January 2, 2015*, <https://www.dni.gov/files/documents/ICD/ICD%20203%20Analytic%20Standards.pdf>.
422. Blasch et al., "Artificial Intelligence Strategies for National Security and Safety Standards."
423. Weinbaum and Shanahan, "Intelligence in a Data-Driven Age."
424. Vogel et al., "The Impact of AI on Intelligence Analysis."
425. Brantly, "When Everything Becomes Intelligence."
426. Sanclemente, "Reliability."
427. Yeung, "When AI Misjudgment Is Not an Accident."
428. Vogel et al., "The Impact of AI on Intelligence Analysis."
429. Mitchell et al., "The Future of Intelligence Analysis."
430. Ganor, "Artificial or Human."
431. Vogel, "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce."
432. Harding, "Move Over JARVIS, Meet OSCAR."
433. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
434. Sanclemente, "Reliability."
435. Sanclemente, "Reliability."
436. Yeung, "When AI Misjudgment Is Not an Accident."
437. Sanclemente, "Reliability."
438. Phillips and Pohl, "Countering Intelligence Algorithms."
439. Jani, "The Promise and Prejudice of Big Data in Intelligence Community."
440. Ish, Ettinger, and Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis."
441. Zegart, "9/11: Look Back and Learn."
442. Blais and Jungdahl, "Artificial Intelligence in a Human Intelligence World."
443. Harding, "Move Over JARVIS, Meet OSCAR."
444. Stone, "Artificial Intelligence in the Intelligence Community: Culture is Critical."
445. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
446. Harding, "Move Over JARVIS, Meet OSCAR."
447. Stone, "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition."
448. Stone, "Artificial Intelligence in the Intelligence Community: Culture is Critical."
449. Ish, Ettinger, and Ferris, "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis."
450. Tarraf et al., "The Department of Defense Posture for Artificial Intelligence."
451. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
452. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
453. Flournoy, Haines, and Chetifz, "Building Trust through Testing."
454. Haines and O'Sullivan, "Maintaining the Intelligence Edge."
455. Ocasio, "Towards an Attention-Based View of the Firm," and Ocasio, "Attention to Attention."
456. Steve W.J. Kozlowski and Katherine J. Klein, "A Multilevel Approach to Theory and Research in Organizations: Contextual, Temporal, and Emergent Process," in Katherine J. Klein and Steve W.J. Kozlowski (eds.), *Multilevel Theory, Research, and Methods in Organizations: Foundations, Extensions, and New Directions*. (San Francisco, CA: Jossey-Bass, 2010), 3-90.

- 457. Charles M. Farkas and Suzy Wetlaufer, "The Ways Chief Executive Officers Lead." *Harvard Business Review*, 74, no. 3 (1996): 110-122. <https://hbr.org/1996/05/the-ways-chief-executive-officers-lead>
- 458. Farkas and Wetlaufer, "The Ways Chief Executive Officers Lead."
- 459. Adrian Wolfberg, "The President's Daily Brief: Managing the Relationship Between Intelligence and the Policy-maker," *Political Science Quarterly* 132, no. 2 (2017): 225-258. <https://doi.org/10.1002/polq.12616>

References

Entries beginning with an asterisk (*) are the 41 references used in the AI/ML in IC literature review.

- *Alessa, et al. 2021. "Surprise and Suspense: How the Intelligence Community Forgot the Future." *The International Journal of Intelligence, Security, and Public Affairs* 23, no. 3: 310-342. <https://doi.org/10.1080/23800992.2021.2006954>
- *Allen, Greg, and Taniel Chan. 2017. "Artificial Intelligence and National Security." Harvard University, July 2017. <https://www.belfercenter.org/sites/default/files/files/publication/AI%20NatSec%20-%20final.pdf>
- *Blais, J. Rocco, and Adam M. Jungdahl. 2019. "Artificial Intelligence in a Human Intelligence World." *American Intelligence Journal* 36, no. 1: 108-113. <https://www.jstor.org/stable/27066342>
- *Blasch, et al. 2018. "Artificial Intelligence Strategies for National Security and Safety Standards." Paper presented at the AAAI FSS-19, Artificial Intelligence in Government and Public Sector conference, Arlington, VA, November 7-9, 2019. <https://arxiv.org/abs/1911.05727>
- *Brantly, Aaron F. 2018. "When Everything Becomes Intelligence: Machine Learning and the Connected World." *Intelligence and National Security* 33, no. 4: 562-573. <https://doi.org/10.1080/02684527.2018.1452555>
- Briner, Rob B., David Denyer, and Denise M. Rousseau. 2009. "Evidence-Based Management: Concept Cleanup Time?" *Academy of Management Perspectives*, 23, no. 4, 19-32: <https://doi.org/10.5465/amp.23.4.19>
- *Bury, Patrick, and Michael Chertoff. 2020. "New Intelligence Strategies for a New Decade." *The RUSI Journal* 165, no. 4: 42-53. <https://doi.org/10.1080/02684527.2018.1452555>
- Cao, Longbing. 2017. "Data Science: A Comprehensive Overview." *ACM Computing Surveys*, 50, no. 3: 1-42. <https://doi.org/10.1145/3076253>
- Chadwick, Andrew, and James Stanyer. 2022. "Deception as a Bridging Concept in the Study of Disinformation, Misinformation, and Misperceptions: Toward a Holistic Framework." *Communication Theory*, 32, no. 1: 1-24. <https://doi.org/10.1093/ct/qtab019>
- Cherry, Gemma, and Rumona Dickson. 2017. "Defining My Review Question and Identifying Inclusion and Exclusion Criteria." In *Doing a Systematic Review: A Student's Guide*, edited by Angela Boland, Gemma Cheery, and Rumona Dickson, 43-59. Los Angeles, CA: SAGE.
- *DeHenre, Rena. 2021. "Why the Department of Defense Should Create an AI Red Team." *Over the Horizon*, September 7, 2021. <https://othjournal.com/2021/09/07/why-the-department-of-defense-should-create-an-ai-red-team/>
- Fallis, Don. 2014. "The Varieties of Disinformation." In *The Philosophy of Information Quality*, edited by Luciano Floridi and Phyllis Illari, 135-162. New York, NY: Springer.
- Farkas, Charles M., and Suzy Wellaufer. 1996. "The Ways Chief Executive Officers Lead." *Harvard Business Review* 74, no. 3: 110-122. <https://hbr.org/1996/05/the-ways-chief-executive-officers-lead>
- *Flournoy, Michele A., Avril Haines, and Gabrielle Chefitz. 2020. "Building Trust through Testing: Adapting DoD's Test & Evaluation, Validation & Verification (TEVV) Enterprise for Machine Learning Systems, including Deep Learning Systems." *Georgetown University*, October 2020. <https://cset.georgetown.edu/wp-content/uploads/Building-Trust-Through-Testing.pdf>

- *Ganor, Boaz. 2021. "Artificial or Human: A New Era of Counterterrorism Intelligence." *Studies in Conflict & Terrorism*, 44, no. 7: 605-624. <https://doi.org/10.1080/1057610X.2019.1568815>
- Gorman, Christopher. 2022. "Recent Developments in AI and National Security: What You Need to Know." *Lawfare*, March 3, 2022. <https://www.lawfareblog.com/recent-developments-ai-and-national-security-what-you-need-know#>
- *Haines, Avril, and Stephanie O'Sullivan. 2021. "Maintaining the Intelligence Edge: Reimagining and Reinventing Intelligence through Innovation." *Center for Strategic and International Studies*, January 2021. https://csis-website-prod.s3.amazonaws.com/s3fs-public/publication/210113_Intelligence_Edge.pdf
- *Hampel-Arias, Zeigried, and John S. Meyers. 2021. "What AI Can and Cannot Do for the Intelligence Community." *Defense One*, January 5, 2021. <https://www.defenseone.com/ideas/2021/01/what-ai-can-and-cannot-do-intelligence-community/171195/>
- *Harding, Emily. 2022. "Move Over JARVIS, Meet OSCAR." *Center for Strategic and Intelligence Studies*, January 2022. https://csis-website-prod.s3.amazonaws.com/s3fs-public/publication/220119_Harding_MoveOverJARVIS_MeetOSCAR_0.pdf?NqfrbU05ULzzcySzNHB0pTzsNYw3HdFK
- Hawk, Thomas F., and Dale E. Zand. 2014. "Parallel Organization: Policy Formulation, Learning, and Interdivision Integration." *Journal of Applied Behavioral Science*, 50, no. 3: 307-336. <https://doi.org/10.1177%2F0021886313509276>
- *Ish, Daniel, Jared Ettinger, and Christopher Ferris. 2021. "Evaluating the Effectiveness of Artificial Intelligence Systems in Intelligence Analysis." *RAND Corporation*. <https://apps.dtic.mil/sti/pdfs/AD1146077.pdf>
- *Jani, Karen. 2016. "The Promise and Prejudice of Big Data in Intelligence Community." *Georgia Institute of Technology*, <https://arxiv.org/pdf/1610.08629v1.pdf>
- *Katz, Brian. 2020. "The Collection Edge: Harnessing Emerging Technologies for Intelligence Collection." *Center for Strategic and International Studies*. July 1, 2020. <https://www.jstor.org/stable/resrep25236?seq=1>
- Knippen, Jay T., and Thad B. Green. 1997. "Problem Solving." *Journal of Workplace Learning*, 9, no. 3: 98. <http://dx.doi.org/10.1108/13665629710164904>
- Kozlowski, Steve W.J., and Katherine J. Klein. 2010. "A Multilevel Approach to Theory and Research in Organizations: Contextual, Temporal, and Emergent Processes." In *Multilevel Theory, Research, and Methods in Organizations: Foundations, Extensions, and New Directions*, edited by Katherine J. Klein and Steve W.J. Kozlowski, 3-90. San Francisco, CA: Jossey-Bass.
- *Laird, John, Charan Ranganath, and Samuel Gershman. 2019. "Future Directions in Human Machine Teaming Workshop." Report prepared at the Future Directions in Human Machine Teaming Workshop series, Arlington, VA, July 16-17, 2019. <https://basicresearch.defense.gov/Portals/61/Future%20Directions%20in%20Human%20Machine%20Teaming%20Workshop%20report%20%20%28for%20public%20release%29.pdf>
- *McCormick, David H., Charles E. Luftig, and James M. Cunningham. 2020. "Economic Might, National Security, and the Future of American Statecraft." *Texas National Security Review* 3, no. 3: 50-75. https://repositories.lib.utexas.edu/bitstream/handle/2152/83223/04_TNSRVol3Issue3McCormick.pdf?sequence=2
- *McNamara, Laura A. 2012. "Interdisciplinary Research in the National Laboratories." In *Anthropologists in the Security-Scape*, edited by Robert Albrow, George March, Laura A. McNamara, and Monica Schoch-Spana, 87-100. Walnut Creek, NY: Left Coast Press.
- *McNamara, Laura A. 2020. "Innovation Adoption, Technology Acceptance, and Data Science: Why Algorithmic Technologies are so Tricky." Slides presents at the Department of Defense Human Systems Integration Community of Practice Seminar Series, Albuquerque, NM, September 4, 2020.
- *Mitchell, Kwasi, Joe Mariani, Adam Routh, Akash Keyal, and Alex Mirkow. 2019. "The Future of Intelligence Analysis: A Task-Level View of the Impact of Artificial Intelligence on Intelligence Analysis." *Deloitte Insights*. https://www2.deloitte.com/content/dam/insights/us/articles/6306_future-of-intel-analysis/DI_Future-of-intel-analysis.pdf

- *Musser, Micah, and Ashton Garriott. 2021. "Machine Learning and Cybersecurity: Hype and Reality." *Georgetown University Center for Security and Emerging Technology*. June 2021. <https://cset.georgetown.edu/wp-content/uploads/Machine-Learning-and-Cybersecurity.pdf>
- *Phillips, Peter J., and Gabriela Pohl. 2020. "Countering Intelligence Algorithms: Decision Theory, Design Choices & Counter-AI." *The RUSI Journal* 165, no. 7: 22-32. <https://doi.org/10.1080/03071847.2021.1893126>
- Ocasio, William. 1997. "Towards an Attention-Based View of the Firm." *Strategic Management Review* 18: 187-206. [https://doi.org/10.1002/\(SICI\)1097-0266\(199707\)18:1+%3C187::AID-SMJ936%3E3.0.CO;2-K](https://doi.org/10.1002/(SICI)1097-0266(199707)18:1+%3C187::AID-SMJ936%3E3.0.CO;2-K)
- Ocasio, William. 2011. "Attention to Attention." *Organization Science* 22, no. 5: 1286-1296. <https://doi.org/10.1287/orsc.1100.0602>
- Office of the Director of National Intelligence, *Intelligence Community Directive 203, Analytic Standards*, January 2, 2015, <https://www.dni.gov/files/documents/ICD/ICD%20203%20Analytic%20Standards.pdf>.
- *Ogunlana, Sunday O. 2019. "Halting Boko Haram/Islamic State's West African Province Propaganda in Cyberspace with Cybersecurity Technologies." *Journal of Strategic Security* 12, no. 1: 72-106. <https://www.jstor.org/stable/26623078?seq=1>
- Rojon, Celine, Adun Okupe, and Almuth McDowall. 2021. "Utilization and Development of Systematic Reviews in Management Research: What Do We Know and Where Do We Go from Here?" *International Journal of Management Reviews* 23, no. 2: 191-223. <https://doi.org/10.1111/ijmr.12245>
- Rousseau, Denise, M., Joshua Manning, and David Denyer. 2008. "Evidence in Management and Organizational Science: Assembling the Field's Full Weight of Scientific Knowledge Through Synthesis." *Academy of Management Annals*, 2, no. 1: 475-515. <https://doi.org/10.5465/19416520802211651>
- Saldaña, Johnny. 2021. *The Coding Manual for Qualitative Researchers*, 4th edition. Thousand Oaks, CA: SAGE.
- *Sanclemente, Guadys L. 2021. "Reliability: Understanding Cognitive Human Bias in Artificial Intelligence for National Security and Intelligence Analysis." *Security Journal* 1-21, <https://doi.org/10.1057/s41284-021-00324-z>
- *Schmidt, Eric, Bob Work, et al. 2021. "Final Report: National Security Commission on Artificial Intelligence (AI)." *National Security Commission on Artificial Intelligence*, <https://apps.dtic.mil/sti/pdfs/AD1124333.pdf>
- *Shedd, David R. 2020. "The Intelligence Posture America Needs in an Age of Great-Power Competition." *The Heritage Foundation*, November 17, 2020. <https://www.heritage.org/military-strength-topical-essays/2021-essays/the-intelligence-posture-america-needs-age-great-power>
- Smith, et al. 2011. "Methodology in Conducting a Systematic Review of Systematic Reviews of Healthcare Interventions." *BMC Medical Research Methodology*, 11, no. 15: 1-6. <https://bmcmmedresmethodol.biomedcentral.com/track/pdf/10.1186/1471-2288-11-15.pdf>
- *Stone, Corin R. 2021. "Artificial Intelligence in the Intelligence Community: Money is Not Enough." *Just Security*, July 12, 2021. <https://www.justsecurity.org/77354/artificial-intelligence-in-the-intelligence-community-money-is-not-enough/>
- *Stone, Corin R. 2021. "Artificial Intelligence in the Intelligence Community: Culture is Critical." *Just Security*, August 17, 2021. <https://www.justsecurity.org/77783/artificial-intelligence-in-the-intelligence-community-culture-is-critical/>
- *Stone, Corin R. 2021. "Artificial Intelligence in the Intelligence Community: The Tangled Web of Budget and Acquisition." *Just Security*, September 28, 2021. <https://www.justsecurity.org/78362/artificial-intelligence-in-the-intelligence-community-the-tangled-web-of-budget-acquisition/>
- *Stone, Corin R. 2021. "Artificial Intelligence in the Intelligence Community: Know Risk, Know Reward." *Just Security*, October 19, 2021. <https://www.justsecurity.org/78641/artificial-intelligence-in-the-intelligence-community-know-risk-know-reward/>

- *Stone, Corin R. 2021. "Artificial Intelligence in the Intelligence Community: Oversight Must Not Be an Oversight" *Just Security*, November 30, 2021. <https://www.justsecurity.org/79254/artificial-intelligence-in-the-intelligence-community-oversight-must-not-be-an-oversight/>
- *Tarraf, Danielle, William Shelton, Edward Parker, et al. 2019. "The Department of Defense Posture for Artificial Intelligence: Assessment and Recommendations." https://www.rand.org/pubs/research_reports/RR4229.html
- *Tucker, Patrick. 2020. "Spies like AI: The Future of Artificial Intelligence for the US Intelligence Community." *Defense One*, January 27, 2020. <https://www.defenseone.com/technology/2020/01/spies-ai-future-artificial-intelligence-us-intelligence-community/162673/>
- *Vogel, Kathleen M. 2021. "Big Data, AI, Platforms, and the Future of the U.S. Intelligence Workforce: A Research Agenda." *IEEE Technology and Society Magazine* 40, no. 3: 84-92. <https://doi.org/10.1109/MTS.2021.3104384>
- *Vogel, et al. 2021. "The Impact of AI on Intelligence Analysis: Tackling Issues of Collaboration, Algorithmic Transparency, Accountability, and Management." *Intelligence and National Security* 36, no. 6: 827-848. <https://doi.org/10.1080/02684527.2021.1946952>
- *Weinbaum, Cortney, and John N.T. Shanahan. 2018. "Intelligence in a Data-Driven Age." *Joint Forces Quarterly* 90, no. 3: 4-9. https://ndupress.ndu.edu/Portals/68/Documents/jfq/jfq-90/jfq-90_4-9_Weinbaum-Shanahan.pdf
- Weinbaum, Cortney, Bradley Knopp, Soo Kim, and Yuliya Shokh. 2022. "Options for Strengthening All-Source Intelligence." *RAND Corporation*, <https://apps.dtic.mil/sti/pdfs/AD1161438.pdf>
- *Wilner, Alex, Casey Babb, and Jessica Davis. 2021. "Four things to consider on the Future of AI-enabled Deterrence." *Lawfare*. Last modified July 25, 2021, 10:01. <https://www.lawfareblog.com/four-things-consider-future-ai-enabled-deterrence#>
- Wolfberg, Adrian. 2017. "The President's Daily Brief: Managing the Relationship between Intelligence and the Policy-maker." *Political Science Quarterly* 132, no. 2: 225-258. <https://doi.org/10.1002/polq.12616>
- Xiao, Yu, and Maria Watson. 2019. "Guidance on Conducting a Systematic Literature Review." *Journal of Planning Education and Research*. 39, no. 1: 93-112. <https://doi.org/10.1177%2F0739456X17723971>
- *Yeung, Douglas. 2018. "When AI Misjudgment Is Not an Accident." *Scientific American*, October, 19, 2018. <https://blogs.scientificamerican.com/observations/when-ai-misjudgment-is-not-an-accident/>
- *Yeung et al. 2021. "Identifying Systemic Bias in The Acquisition of Machine Learning Decision Aids for Law Enforcement Applications." *RAND Corporation*, <https://www.rand.org/pubs/perspectives/PEA862-1.html>
- *Zegart, Amy. 2020. "9/11: Look Back and Learn." *Hoover Digest*, January 20, 2020. <https://www.hoover.org/research/911-look-back-and-learn>
- *Zegart, Amy. 2022. "American Spy Agencies Are Struggling in the Age of Data." *Wired*, February 2, 2022. <https://www.wired.com/story/spies-algorithms-artificial-intelligence-cybersecurity-data/#:~:text=In%20short%2C%20data%20volume%20and,struggling%20to%20adapt%20to%20it.>

