

Article

Exploring the Interdependence Theory of Complementarity with Case Studies. Autonomous Human–Machine Teams (A-HMTs)

William F. Lawless

Department of Mathematics, Sciences and Technology, School of Arts & Sciences, Paine College,
Augusta, GA 30901, USA; wlawless@paine.edu

Abstract: Rational models of human behavior aim to predict, possibly control, humans. There are two primary models, the cognitive model that treats behavior as implicit, and the behavioral model that treats beliefs as implicit. The cognitive model reigned supreme until reproducibility issues arose, including Axelrod's prediction that cooperation produces the best outcomes for societies. In contrast, by dismissing the value of beliefs, predictions of behavior improved dramatically, but only in situations where beliefs were suppressed, unimportant, or in low risk, highly certain environments, e.g., enforced cooperation. Moreover, rational models lack supporting evidence for their mathematical predictions, impeding generalizations to artificial intelligence (AI). Moreover, rational models cannot scale to teams or systems, which is another flaw. However, the rational models fail in the presence of uncertainty or conflict, their fatal flaw. These shortcomings leave rational models ill-prepared to assist the technical revolution posed by autonomous human–machine teams (A-HMTs) or autonomous systems. For A-HMT teams, we have developed the interdependence theory of complementarity, largely overlooked because of the bewilderment interdependence causes in the laboratory. Where the rational model fails in the face of uncertainty or conflict, interdependence theory thrives. The best human science teams are fully interdependent; intelligence has been located in the interdependent interactions of teammates, and interdependence is quantum-like. We have reported in the past that, facing uncertainty, human debate exploits the interdependent bistable views of reality in tradeoffs seeking the best path forward. Explaining uncertain contexts, which no single agent can determine alone, necessitates that members of A-HMTs express their actions in causal terms, however imperfectly. Our purpose in this paper is to review our two newest discoveries here, both of which generalize and scale, first, following new theory to separate entropy production from structure and performance, and second, discovering that the informatics of vulnerability generated during competition propels evolution, invisible to the theories and practices of cooperation.



Citation: Lawless, W.F. Exploring the Interdependence Theory of Complementarity with Case Studies. Autonomous Human–Machine Teams (A-HMTs). *Informatics* **2021**, *8*, 14. <https://doi.org/10.3390/informatics8010014>

Academic Editor: Antony Bryant

Received: 20 January 2021

Accepted: 18 February 2021

Published: 26 February 2021

Publisher's Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright: © 2021 by the author. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Keywords: interdependence; complementarity; bistability; uncertainty and incompleteness; non-factorable information and tradeoffs; autonomous teams (Level 5 vehicles)

1. Introduction

As part of the background, we revisit issues previously identified regarding teams, organizations, and social systems in preparation for the dramatic arrival of autonomous human–machine teams (A-HMTs) in the military (e.g., hypersonic missiles), science (e.g., transportation systems; medical systems), and society (for a review, see reference [1]). Autonomous machines imply that artificial intelligence (AI, including machine learning, or ML) is needed to produce autonomy along with the extraordinary technological changes that have been occurring recently (e.g., see the Rand report on AI, in [2]), including the possibility of autonomous military weapon systems ([3]; e.g., limited autonomy is already planned to be given to swarms; [4]). Our purpose is to provide a review of what we have accomplished, and a roadmap for further explorations of interdependence and complementarity.

To achieve efficiency, as Feynman [5] warned about classical computers attempting to model quantum systems, AI should not be applied ad hoc; to operate A-HMTs, we must first have a viable theory of interdependence and a sufficient mathematical model of an autonomous human–machine system that can account for predictions and observations under uncertainty and conflict, as one of the first steps to achieve effectiveness, then efficiency. We extend theory to guide research with a model of the interdependence theory of complementarity that incorporates the outcomes of debate, measured by entropy, and a review of our recently discovered extension to vulnerability in organizations.

Background. The 2018 US National Defense Strategy addressed the challenges faced by competition with nuclear powers China and Russia, rogue states Iran and North Korea, a weakening international order, and technological changes offering faster-than-human decisions ([6,7]; also, see the Rand report in [8]). To confront these challenges, among its options, the US Department of Defense (DoD) wants to develop artificial intelligence (AI), autonomy, and robots. Here, we further develop new theory for the application of AI to autonomous human–machine teams (A-HMTs) as part of a system of competing teams. If successful, the theory will be available for the metrics of team performance, managing autonomy, and scaling to organizations and alliances like NATO [9], all necessary to defend a team, corporation, or nation in, for example, a multipolar nuclear world and against hypersonic weapons (e.g., [10]).

Groups have been studied for over a century. Lewin [11] began the scientific study of groups around the social effects of interdependence, observing a “whole is greater than the sum of its parts” (similarly, in Systems Engineering; in [12]). Lewin’s student, Kelley [13], pursued static interdependence with game theory, but, unable to disentangle its effects on player preferences in actual games, Jones (p. 33, [14]) labeled it as bewildering in the laboratory; the predicted value of cooperation, established by games in the laboratory (pp. 7–8, [15]), has not been validated in the real world ([16]; viz., compare p. 413 with 422). Despite its flaws (e.g., reproducibility; in [17]), the rational cognitive model, based on individual consistency, reigns in social science, AI, and military research (e.g., a combat pilot’s Observe-Orient-Decision-Action, or OODA loop; in [18]). It continues to promote cooperation; e.g., in their review of social interdependence theory, Hare and Woods [19] posit anti-Darwinian and anti-competition views in their fulsome support of cooperation, but then, they are unable to generalize, to scale, or to predict outcomes in the face of conflict or uncertainty.

We have countered with: first, the study of groups often uses independent individuals as the unit of analysis instead of teams (p. 235, [20]); second, however, an individual alone cannot determine context nor resolve uncertainty [1]; and third, if teammates assume orthogonal roles, the lack of correlations found among self-reports accounts for the experimental failure of role complementarity [21]; i.e., it has long been proposed that for human mating, opposites in orthogonal roles should be more attracted to each other and should best fit together; however, the lack of correlations has not supported this prediction (p. 207, [22]). Nonetheless, along with the lack of experimental support, by definition, subjective views collected from orthogonal role players should never have been expected to be correlated [23].

We arbitrarily construe the elements of the interdependence of complementarity to include bistability (actor–observer effects, in [24]; two sides to every story; two opposed tribes; individual or group member); uncertainty from measurement; and non-factorable information [21]. Examples: First, the performance of a team of individuals can be greater than the same members who are operating independently instead ([25,26]). Second, humans explore uncertainty with tradeoffs [27]. Third, non-factorable information indicates that information is subadditive, for example, found with economies of scope (e.g., one train can carry both passengers and freight more efficiently than two trains); with risk in an investment portfolio being less than each investment separately [28]; and, with mergers that decrease risk ([29]; e.g., the United Technologies and Raytheon merger in 2019; in [30]). These findings, agreeing with human team research [31], indirectly suggest that a team’s

entropy production is separated into structural and operational streams; that is, the more stable and well-fitted is the structure of a team, the less entropy it produces [23], allowing more free energy directed, say, to explore and to find solutions to problems; e.g., patents.

Adding support for the formulation of our interdependence theory of complementarity, we have established interdependent effects fundamental to a theory of autonomy: Optimum teams operate at maximum interdependence [26]; employee redundancy impedes teams and increases corruption [23]; team intelligence is critical to producing maximum entropy (MEP, in [32]; e.g., we have found that the search to develop patents in the Middle East North African countries, including Israel, depends on the average level of education in a nation; in [23]); and for the public debates offered to persuade a majority in support of an action [21]. In contrast are the weaker decisions by single individuals (authoritarians), by those seeking consensus to control a group, reflecting minority control, or by those that include emotion (e.g., the emotion sometimes associated with divorce or business breakups). Characterizing another failure of the rational model applied to decision-making, the premise of the book on teaching human forecasters about how to become “Superforecasters,” by Tetlock and Gardiner [33], was its inability to predict the Brexit.

2. Materials and Methods

Our goal was to establish whether a theory of interdependence can successfully build and operate A-HMTs effectively, safely, and ethically. Based on theory and case studies, we extend our previous results to the spontaneous debates that arise automatically as a team or system explores its choices by making tradeoffs before an audience to maximize its production of entropy (MEP) as it attempts to solve a targeted problem. This setting often occurs in highly competitive, conflictual, or uncertain situations.

Hypothesis 1 (H1). *First, we hypothesized that a cohesive team with interdependently orthogonal roles will produce uncorrelated information (e.g., the invalidity of implicit attitudes; see [34]), precluding rational decisions, but by allowing a team's fit to become cohesive, produces a subadditive effect (with S as entropy):*

$$S_{A,B} \leq S_A + S_B \quad (1)$$

Equation (1) is reflected by Lewin's ([1]; see also [35]) claim which we reiterate that the sum of the elements can be less than a whole. When the best military teams can be surprised by poorly trained locals, a cohesive team is more adaptable to surprise [36], to unexpected shocks [21]: and more likely to build trust [1]. However, the whole can become less than the sum of its parts too (e.g., a divorce; a fight internal to an organization; or an operator on a team who fails to perform). Humans must learn to trust machines, but, once trained as part of a human-machine team, these trained machines know how the humans are expected to behave. Then, as a duty to protect the members of a team, given the authority to intervene when a human operator becomes dysfunctional performing its role in a team, the machine might be able to save lives; e.g., an airliner preventing a copilot from committing suicide, or a fighter plane placing itself in a safe state if its pilot passes out from excessive g-forces [37].

Until recently, human-centered design (HCD) dominated research for twenty years [38] in making designs, like a self-driving a car and more complex ones as well. As justification, Cooley [39] claimed that humans must come before machines, no matter the value a machine might provide. However, at one time, HCD was once considered harmful (e.g., [40]), like autonomous systems are considered today (e.g., [41]).

Autonomy elevates the risks to humans. Human observers “in-the-loop” serve to provide a check on autonomous systems; in contrast, “on-the-loop” observations by humans of autonomous machines carry the greater risks to the public. The editors of the *New York Times* [42] wrote about their concern that autonomous military systems could be overtaken and controlled by adversaries wittingly or unwittingly. Paraphrasing the Editors who

quoted UN Secretary General, Antonio Guterres, autonomous military machines should be forbidden by international law from taking human lives without the oversight provided by humans. The editors went on to conclude that humans should never assign life and death decisions to autonomous machines. But, and just the opposite conclusion, because human error is the cause of most accidents [37], autonomous machines might save lives.

Lethal Autonomous Weapon Systems (LAWS; in [3]), however, may independently identify and engage a target without human control. Rapid changes in technology suggest that current rules for LAWS are already out of date [43]. For autonomous systems to exist and operate effectively and efficiently, and safely and ethically, even in the context of command and control, effective deterrence, and enhanced security in a multipolar nuclear world, we must be able to design trustworthy A-HMTs that operate rapidly and cohesively in multi-domain formations. As a caution, multi-domain operations rely on convergence processes [44] which may work in highly certain environments, but not when uncertainty prevails; worse, we argue, by dismissing an opponent's reasoning, thereby adding incompleteness to a context, convergence processes may increase uncertainty [21].

Hypothesis 2 (H2). *As our second hypothesis, we considered that the role of debate for humans is to test in tradeoffs how a team (system) can increase its production of entropy to a maximum (MEP) as attempts are made to compete in conflictual or uncertain situations.*

A team has a plan to operate in uncertain environments. When a plan performs well, even for multi-domain operations, its plan should be executed. When it faces an uncertain context [1], guided by interdependence theory, we recommend debate on tradeoffs in the search for the best decision; e.g., to explore their options, U.S. Presidents Washington and Eisenhower both encouraged strong dissent (respectively, [45–47]).

We found in our review of the literature [48] that “human behavior and actions remains largely unpredictable.” Physical network scientists, like Barabasi [49], and modern game theorists [50] disagree; in addition, physical network scientists not only want to be able to predict behavior, but also to control behavior [51]. Rejecting the cognitive model dramatically improved their predictability of behavior, but only when beliefs are suppressed, in low risk environments, or highly certain environments. Inversely, the cognitive model diminishes the value of behavior [52]. These models, however, fail in the presence of uncertainty [53] or conflict [54], exactly where the interdependence theory of complementarity thrives [23]. Confronting conflict or uncertainty, the alternative views of reality are spontaneously tested by debating the different interpretations that arise whenever humans search for an optimal path so that they can proceed to their goal. This common human reaction to the unknown means that for machines to partner with humans, machines must be able to discuss their situation, plans, decisions and actions with their human teammates in causal terms that both can understand albeit imperfectly at this time ([55,56]).

3. Results of Case Studies

Case study 1: An Uber car was involved in the fatal death of a pedestrian in 2018 ([57,58]). Its machine-learning correctly acquired the pedestrian 6 s early compared to its human operator at 1 s, but the car's access to its brakes had been disabled to improve the car's ride. However, establishing the machine's lack of interdependence with its human operator, and unlike a good team player, the car failed to alert its human operator when it first acquired an object on the road, an action made more likely in the future with AI, but less likely with machine learning.

Case 2. In 2017–2018, teams composed of humans and robots provided synergism to BMW's operations (positive interdependence) that allowed BMW to grow its car production with more robots and workers [59]. At the same time, Tesla's all-robot operations were in trouble. Despite its quarterly quota of only 5000 cars, Tesla struggled to meet its goal [60]. This adverse situation was not solved by its managers, by its human operators, nor by Tesla's robots. To make its quota of 5000 units per quarter, Tesla removed and replaced many of its robots with humans. Later analysis discovered that neither its humans or its

robots could see that the Tesla robots were dysfunctional in the positions that had to assume while they were assembling cars. After improving the vision of its robots [61], Tesla's car production soared well beyond its quarterly quota to over 80,000 units in 2018 [62] and 2019 [63].

Case study 3. Two collisions occurred at sea in 2017, first by the USS Fitzgerald which killed seven sailors, and kept the Fitzgerald, a guided-missile destroyer, out of service for almost two years [64]. The second collision was by the USS McCain. The latter suggests that an interdependent A-HMT may make operations safer in the Naval Fleet. From the US National Transportation Safety Board (NTSB; p. viii, [65]), the destroyer John S McCain and the tanker Alnic MC collision was attributed in part to inadequate bridge operating procedures. The John S McCain's bridge team lost situational awareness (context) and failed to follow its loss of steering emergency procedures, particularly to alert nearby traffic of their perceived loss of steering. Contributing to the accident was the operation of the steering system in backup manual mode, which allowed for an unintentional, unilateral transfer of steering control. These problems may have been prevented by an aware machine teammate (i.e., the ship) interdependent with the bridge's crew, a machine given the authority to intervene to prevent a collision and save a ship.

Case study 4: The decision processes at the U.S. Department of Energy's (DOE) Citizens Advisory Board (CAB) at DOE's Hanford site in the State of Washington and DOE's Savannah River Site (SRS) in South Carolina are compared, along with the results of those decisions. Since its founding, Hanford's CAB has forcibly sought to reach consensus in all of its decisions (also known as control by a minority, because a few members can become a barrier to action by blocking a majority at any time), a process that has promoted conflict among CAB members [66] and impeded the action by DOE's Hanford to close any of its radioactive high-level waste (HLW) tanks [67]. SRS closed its first two HLW tanks under regulatory oversight in the world in 1997, and several since. Almost 25 years after SRS began closing its HLW tanks, recently, Hanford entered into a contract to begin to close its first HLW tank [68]. In contrast to the barriers imposed by the consensus-seeking rules of Hanford's CAB, majority-rules used by the DOE SRS in South Carolina have accelerated the cleanup at SRS motivated by its CAB's support. Moreover, the majority rules at SRS have also promoted collegial relations among its citizens and DOE site managers that has rapidly and safely cleaned up SRS ([66]; see also [23]).

4. Discussion of the Results

Case study 1. Human teams are autonomous. For human-machine teams and systems, what does autonomy require? In the first case study, from a human-machine team's perspective, clearly, the Uber car was an inept team player [1]. Facing uncertain situations, the NTSB report confirmed that alone, a single human, a single machine or a single team is unable to determine a rapidly changing context (e.g., driving at night while approaching an unknown object in the road); indeed, as we have argued [1], resolving an uncertain context requires at a minimum a bistable state of disagreement shared between two or more interdependent agents to adapt to rapid changes in context (e.g., by, say, sharing different interpretations of reality), and, overall, to operate safely and ethically autonomous human-machine systems. We also know from Cummings [26] that the best science teams are fully interdependent. Cooke [69] attributes the intelligence of a team in the interdependent interactions among its teammates (when the information gleaned is shared). To reduce uncertainty in an autonomous team requires that teammates, whether human or machine (robot, self-driving car, etc.), can share their interpretations of reality, however imperfectly determined, in causal terms ([55,56]). In the case of the Uber car, the self-driving vehicle failed to share the information it had at the time of the fatal accident.

Case studies 2–3. That the intelligent machine in case study 1 stood mute when its human teammate needed to know their change in the context was also evident in case studies 2 and 3. In these two cases, fully equipped machines stood mute or passive while the humans struggled to determine the quickly changing context that they had entered. In

contrast to the struggling humans, the intelligent, very capable machines stood passively quiet, and unfortunately by design, unable to help.

Case study 4. Here, we go into more detail of the decision–action processes at SRS. After SRS had closed two of its HLW tanks and as it began to close two more, DOE was sued, lost the suit which forced it to stop closing HLW tanks, but then, was rescued by Congress. Congress allowed DOE to resume its HLW tank closures if, and only if, the U.S. Nuclear Regulatory Commission (NRC) agreed. Seven years later, NRC had not yet agreed to close the next two HLW tanks. We recently designed an equation (Equation (2), discussed later) to model the seven years-long debate between DOE and NRC over the closure of DOE’s HLW tanks at its SRS site as the situation stood in 2011 [23]. Operating as a quasi-Nash equilibrium between relatively equal opponents, but where neither side is fully able to grasp reality, we located opposing views on an imaginary y -axis (points 1, 2 in Figure 1). The two competing teams of DOE and NRC acted like a capacitor, each side storing then releasing orthogonal information to induce an audience to process both sides as it oscillates until a decision is rendered. Under consensus-seeking rules (designed to promote cooperation, but, counterintuitively, by increasing frustration, they replace the positive aspects of debate with conflict; in [66]), the decision is spread among the audience members; the problem with consensus-seeking arises in that any one is permitted to block an action, modeled in Figure 1 by making points 1 and 2 very stable, but, by not processing information, it enables the minority control preferred by autocrats. Compared to the rules for seeking consensus (viz., Hanford’s CAB), the audience participating in majority rule (viz., SRS’s Board) was able to bring fierce pressure against DOE to take action by closing more tanks, modeled as resistance (point 6 in Figure 1).

As another example supporting this discussion of the results, in a study by the European Council (from p. 29, in [70]), “The requirement for consensus in the European Council often holds policy-making hostage to national interests in areas which Council could and should decide by a qualified majority.”

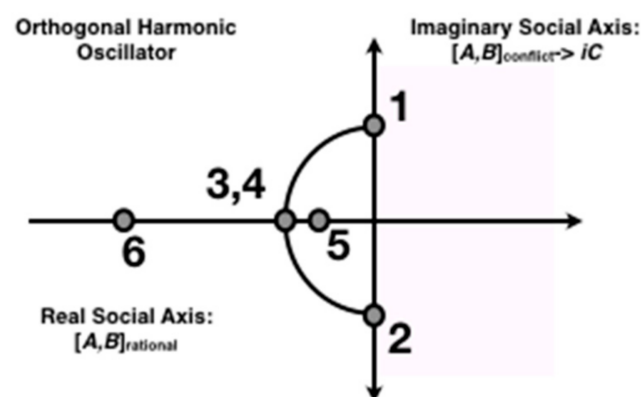


Figure 1. (From [71], submitted). The points (1 and 2) on the imaginary y -axis represent beliefs posed as choices to be debated for action. If no consensus for an agreement to take *action* is reached, no action occurs (e.g., the 7-year stalemate in the debate between DOE and NRC). Points 3 and 4 reflect a compromise choice for action on the x -axis (physical reality); points 5 and 6 represent resistance to the debaters that their audience has decided on the action it wants them to execute, with point 6 reflecting the strongest social feedback by the audience.

The movement towards autonomy begins with a return to the whole. We apply two ideas from traditional social science to our concept of autonomous systems: the separation between structure [23] and function [31]. With this separation in mind, we consider whether the separation of the structure of autonomous human–machine participants for a team in their interactions during their performance affords their team an advantage.

We reconsider Equation (1). In 1935 (p. 555, [72]), Schrödinger wrote about quantum theory by describing entanglement: “the best possible knowledge of a *whole* does not

necessarily include the best possible knowledge of all its *parts*, even though they may be entirely separate and therefore virtually capable of being ‘best possibly known’ ... The lack of knowledge is by no means due to the interaction being insufficiently known ... it is due to the interaction itself.” Returning to Lewin (p. 146, [11]), the founder of Social Psychology, Lewin wrote that the “whole is greater than the sum of its parts.” Moreover, from the Systems Engineering Handbook [12], “A System is a set of elements in interaction” [73] where systems “often exhibit emergence, behavior which is meaningful only when attributed to the whole, not to its parts” [74]. Now, applying Schrödinger to Lewin’s and Systems Engineering’s interpretation of the “whole” begins to account for the whole being greater than the sum of its parts.

There is more to be gained from Schrödinger (again on p. 555, [72]): “Attention has recently been called to the obvious but very disconcerting fact that even though we restrict the disentangling measurements to one system, the representative obtained for the other system is by no means independent of the particular choice of observations which we select for that purpose and which by the way are entirely arbitrary.”

If parts of a whole are not independent (Equation (1)), does a state of interdependence among the orthogonal, complementary parts of a team confer an advantage to the whole [1]? Justifying Equation (1), an answer comes from the science of teams: Compared to a collection of the same but independent individuals, the members of a team when interdependent are significantly more productive ([25,26]). An autonomous whole, then, means a loss of independence among its parts; i.e., the independent parts must fit together into a “structural” whole. Thus, for the whole of an autonomous system to be greater than the sum of the whole’s individual parts, the structure and function of the whole must be treated separately ([31]; see also [23]).

As the parts of a whole become well-fitted together, supporting Equation (1), in the limit, its degrees of freedom (*dof*) reduce, thereby reducing the entropy produced as shown in Equation (2):

$$S_{A,B} = \lim_{dof \rightarrow 1} \log(dof) = 0 \quad (2)$$

The reader may ask, “When the whole is greater than the sum of its parts, why is that not superadditive?” Superadditivity is a rational approach to summing the parts of a whole in which case, entropy generation equals or exceeds inputs (e.g., a system with minimal stability, nonlinear process, or internal irreversibility). Subadditive refers to a reduction in entropy (e.g., life). In our case, where the whole is greater than the sum of its contributing parts, subadditivity is counterintuitive. It demands a team, composed of individual parts operating interdependently as a whole. It works best when its parts are in orthogonal roles, allowing the parts to fit together, each part depending on the whole to operate as a close-knit team (e.g., a small repair business with a clerk, a mechanic and an accountant).

Future research. Next, we plan to model the potential to change free energy, α , per unit of entropy ($d\alpha/ds$), which gives us Equation (3):

$$d\alpha/dt = d\alpha/ds * ds/dt = \tau_A \omega_A = \tau_B \omega_B \quad (3)$$

Equation (3) offers the possibility of a metric, which we devise to reflect a decision advantage, DA :

$$DA = \tau_{output} / \tau_{input} \quad (4)$$

Equations (3) and (4) measure how rapidly the information from one team is circulating to another and back again during a debate, our second recent discovery with the interdependence theory of complementary. However, DA has implications for the application of theory to states of teams during a competition.

Namely, decision advantage has led us to a reconceptualization of competition between equal competitors as a means to seek an advantage by using the forces of competition to discover the information about where an opponent is vulnerable. For example, Boeing’s recent failure to compete against Europe’s Airbus [75] came about from its internal decisions which left it vulnerable after its two 737-Max airliner crashes. The same is true,

however, of multiple other firms, including the HNA Group's unmanageable debt from its frenzy of deal making which left HNA vulnerable to the unexpected collapse of commercial air travel during the pandemic of 2020 and led to its undoing [76]: "From 2015 to 2017, as the Chinese conglomerate HNA Group Co. collected stakes in overseas hotels, financial institutions and office towers in a \$40 billion deal-making frenzy ... [which led to] the downfall of one of China's most acquisitive companies ... HNA found itself in a deep hole because of its wild growth and poor past decisions ... The company grew out of Hainan Airlines ... in the early 1990s. ... HNA piled on debt as it made acquisitions ... In February, when the coronavirus ravaged the airline industry, HNA was effectively taken over by the Hainan government ... to return HNA to financial stability so that it can continue to operate as China's fourth-largest aviation company".

Decision advantage is also implicated in suggesting that mergers are motivated to occur among top competitors to maintain or to increase their competitiveness, or to exploit a weakness in a competitor. The *Wall Street Journal* wrote about Intel's self-inflicted problems as it has struggled with its competitor Advanced Micro Devices (AMD), providing an example of how AMD, the smaller firm, is advancing while the larger firm, Intel, is not [77]: "Advanced Micro Devices Inc. plans to buy rival chip maker Xilinx Inc. in a \$35 billion deal, adding momentum to the consolidation of the semiconductor industry that has only accelerated during the pandemic. AMD and Xilinx on Tuesday said the companies reached an all-stock deal that would significantly expand their product range and markets and deliver a financial boost immediately on closing. ... AMD is enjoying a surge at a time when Intel has been struggling ... stung by problems that analysts believe could help competitors like AMD advance".

Decision advantage can be used by a firm against itself to expose and address its own vulnerabilities and to improve its ability to compete, as in the case with Salesforce in its competition with Microsoft. From *Bloomberg News* [78]: "Salesforce.com Inc. agreed to buy Slack Technologies Inc. ... giving the corporate software giant a popular workplace-communications platform in one of the biggest technology deals of the year. ... The Slack deal would give Salesforce, the leader in programs for managing customer relationships, another angle of attack against Microsoft Corp., which has itself become a major force in internet-based computing. Microsoft's Teams product, which offers a workplace chatroom, automation tools and videoconference hosting, is a top rival to Slack. ... *Bloomberg News* and other publications reported that companies including Amazon.com Inc., Microsoft and Alphabet Inc.'s Google expressed interest in buying Slack at various times when it was still private. ... Salesforce ownership will mark a new era for Slack, a tech upstart with the lofty goal of trying to replace the need for business emails. The cloud-software giant may be able to sell Slack's chatroom product to existing customers around the world, making it even more popular. Slack said in March that it had reached 12.5 million users who were simultaneously connected on its platform, which has grown more essential while corporate employees work from home during the coronavirus pandemic".

5. Discussion

The interdependence theory of complementarity guides us to conclude that intelligence in the interactions of teammates requires that teammates must be able to converse in a causal language that all teammates in an autonomous system understand; viz., to obtain intelligent interactions leads a team to chose teammates who fit together and cohere best. As the parts become a whole in the limit ([1,21]), the informatics of the entropy generated by an autonomous team's or system's structure must be at a minimum to characterize the well-fitted team and to allow the intelligence in the team's interactions to maximize its performance (maximum entropy production, or MEP; in [79]); e.g., by overcoming the obstacles a team faces (e.g., [32]); by exploring a potential solution space for a patent [21]; or by merging with another firm to reduce a system's vulnerability (e.g., [80]) (Huntington Ingalls Industries has purchased a company focused on autonomous systems [80]). In autonomous systems, characterizing vulnerability in the structure of a team, system, or

an opponent was the job in case study 1 that Uber failed to perform in its safety analyses; instead, it became the job that NTSB performed for Uber in the aftermath of Uber's fatal accident. Moreover, the Uber self-driving car and its operator never became a team, remaining instead as independent parts equal to a whole; nor did the Uber car as a teammate recognize that its partner, the operator, had become negligent in her duties and then take the action needed to save itself and its team [37].

The same said for case study 1 can be said about case studies 2 and 3. This means that for machines to become a good teammate requires that they are able to communicate with their human counterparts. It does no teammate or team any good when a multimillion dollar ship sits passively by as the ship collides with another ship. especially when the machine (i.e., the ship) likely could have prevented the accident [37].

In sum, interdependence is critical to the mathematical selection, function, and characterization of an aggregation formed into an intelligent, well-performing unit, achieving MEP in a tradeoff with structure, like in the manner of focusing a telescope. Unfortunately, interdependence also tells us, first, that since each person or machine must be selected in a neutral trial-and-error process [81], the best teams cannot be replicated; and, second, neutrality means that the information for a successful, well-fitted team cannot be obtained in static tests but is only available from the dynamic information afforded by the competitive situations able to stress a team's structure as it performs its functions autonomously; i.e., not every good idea for a new structure succeeds in reality (e.g., a proposed health venture became "unwieldy," in [82]). This conclusion runs contrary to matching theory (e.g., [83]) and to social interdependence theory [19]. However, it holds in the face of uncertainty and conflict for autonomous systems (references [53,54], respectively).

We close with a speculation about the noise and turmoil attending competition, for which few witnesses claim to like (viz., [19]). Entanglement and tunneling have been established as fundamental to cellular activity (McFadden and Al-Khalili [84]). In their review of quantum decoherence, unexpectedly for biological systems, McFadden and Al-Khalili suggest that environmental fluctuations occurring at biologically relevant length and time scales appear to induce quantum coherence, not decoherence; for our research, their review indicates that the quantum likeness that we have posited in the interdependence of complementarity is more likely to occur under the noise and strife of competition that generates the information needed to guide competitive actions rather than the cooperation which consumes or hides information from observers [21]. Instead of supporting the anti-Darwinianism of Hare and Woods [19], the noise and tumult associated with competition motivates the evolution of teams, organizations, and nations.

6. Conclusions

Confronted by conflict or by uncertainty, the rational model fails; instead, the interdependence of complementarity theory succeeds where the rational model does not. In contrast to a rational approach of behavior or cognition, the interdependence theory of complementarity has so far not only successfully passed all challenges, but also, and more importantly, guided us to new discoveries. With case studies, we have advanced the interdependence theory of complementarity with a model of debate, a difficult problem. By the end of this phase of our research, the goal is to further test our theory of interdependence for A-HMTs, to improve the debate model, and to craft performance metrics.

In contrast to traditional social psychological concepts that focus on biases, many of which have failed to be validated (e.g., implicit racism; self-esteem; etc.), the successes of interdependence theory derive from focusing on the factors that increase or decrease the effectiveness of a team; e.g., superfluous redundancy in a team or organization; overcapacity in a market suggesting the need to consolidate; and the need for intelligent interactions to guide a team's decisions as it faces uncertainty or conflict.

The interdependence theory of complementary accounts for why an alternative or competing view of reality helps both interpreters of reality to grasp enough of reality to navigate through an unknown situation, the evidence allowing us to argue that resolving

uncertainty requires a theory of interdependence to construct contexts effectively, efficiently, and safely for autonomous human–machine teams and systems [1].

With interdependence theory, we predicted and replicated that redundancy on a team reduced its ability to be productive [85]. These two successes accounted for the two findings by Cummings [26], first, that the best science teams were highly interdependent from a lack of redundancy on the teams; and, second, that the surprisingly unproductive nature of interdisciplinary science teams were caused by an increase in redundancy as the productive members of a team adapted by working around their unproductive teammates [85]. By extension, we have proposed a generalization of our theory and found that when two equal teams compete against each other, which increases conflict and uncertainty, each team's goal is to find vulnerabilities in the opposing team, characterized by an increase in the losing team's structural entropy production as it struggles to reduce the vulnerability in its structure, illustrating that damage is occurring to its structure. By competing where traditional social science, information theory and artificial intelligence (AI) cannot, the interdependence theory of complementary overthrows key aspects of rational team science.

Autonomous machines one day in the distant future may save lives by replacing humans in safety-critical applications, but not today. We foresee that happening little by little, but not yet. Our contribution with this paper is a review of the difficulty of achieving autonomy with human–machine teams and systems under conditions of uncertainty or conflict. This problem is significantly more difficult than achieving the design of autonomous vehicles, known as Level 5. From Walch [86], “At the ultimate level of autonomy, Level 5 vehicles are completely self-driving and autonomous. The driver does not have to be in control at all during travel, and this vehicle can handle any road condition, type of weather, and no longer bound to geo-fenced locations. A level 5 vehicle will also have emergency features, and safety protocols. A level 5 vehicle is true autonomy and will be able to safely deliver someone to point A to point B. No commercial production of a level 5 vehicle exists, but companies such as Zoox, Google's Waymo, and many others are working towards this goal.” These Level 5 vehicles, however, will for the most part be operating under certainty in the driving conditions they face. The problem we posed in this research article is for autonomous teams and systems solving complex problems in any situation, including during competition, uncertainty or conflict. Level 5 vehicles will not have to engage in debate to resolve the problems they confront. Having said that, however, the fatality caused by the self-driving Uber car illustrated how far from autonomy our science has yet to travel before autonomous teams become reality.

Next, we plan to further explore the foundations of innovation. We found previously that education in a nation was significantly associated with its production of patents [87]. However, this finding was orthogonal to our much earlier finding that an education in air-combat maneuvering had no effect on the performance of combat fighter pilots, both findings that we predicted based on theory. Our theory also depends on the reduction in degrees of freedom for an interdependent team, an idea we have borrowed from Schrödinger (see [72,88]). We plan to further explore its ramifications in the future. We also plan to explore alternative and complementary approaches to different modalities of information [89].

Funding: This research received no external funding.

Institutional Review Board Statement: Not relevant; no human subjects or animals were used, instead, only publicly available data was used and it was cited.

Informed Consent Statement: Not relevant; no human subjects or animals were used, instead, only publicly available data was used and it was cited.

Data Availability Statement: Not relevant; data was not collected nor analyzed. All data referred to in this study was publicly available data and it was cited.

Acknowledgments: This work was completed by the first author alone and without funding from any source. During the last summer in 2020, however, the author was funded by the US Office of Naval Research for summer faculty research, not specifically for the research that led to this manuscript. The author acknowledges that some of the ideas in this paper may have derived from that time.

Conflicts of Interest: The author declares no conflict of interest.

References

1. Lawless, W.F.; Mittu, R.; Sofge, D.A.; Hiatt, L. Editorial (Introduction to the Special Issue), Artificial intelligence (AI), autonomy and human-machine teams: Interdependence, context and explainable AI. *AI Mag.* **2019**, *40*, 5–13. [CrossRef]
2. Tarraf, D.C.; Shelton, W.; Parker, E.; Alkire, B.; Carew, D.G.; Grana, J.; Levedahl, A.; Leveille, J.; Mondschein, J.; Ryseff, J.; et al. The Department of Defense Posture for Artificial Intelligence. Assessment and Recommendations, RAND Corporation. Available online: https://www.rand.org/pubs/research_reports/RR4229.html (accessed on 2 October 2020).
3. Congressional Research Service (CRS). Defense Primer: U.S. Policy on Lethal Autonomous Weapon Systems, Congressional Research Service. Available online: <https://fas.org/sgp/crs/natsec/IF11150.pdf> (accessed on 28 March 2020).
4. Larson, C. Smart Bombs: Military, Defense, National Security and More. Golden Horde: The Air Force’s Plan for Networked Bombs. 1945. Available online: <https://www.19fortyfive.com/2021/01/golden-horde-the-air-forces-plan-for-networked-bombs/> (accessed on 15 January 2021).
5. Feynman, R.P. Quantum mechanical computers. *Found. Phys.* **1986**, *16*, 507–531. [CrossRef]
6. United States Department of Defense. Summary of the 2018 National Defense Strategy of the United States of America. Sharpening the American Military’s Competitive Edge. Available online: <https://dod.defense.gov/Portals/1/Documents/pubs/2018-National-Defense-Strategy-Summary.pdf> (accessed on 4 November 2020).
7. Johnson, J. Artificial Intelligence and the Bomb: Nuclear Command and Control in the Age of the Algorithm. *Modern War Institute*, 7 May 2020.
8. Cohen, R.S.; Chandler, N.; Enron, S.; Frederick, B.; Han, E.; Klein, F.E.; Rhoades, A.L.; Shatz, H.J.; Shokh, Y. Peering into the Crystal Ball. Holistically Assessing the Future of Warfare, Rand Corporation, Document Number: RB-10073-AF. Available online: https://www.rand.org/pubs/research_briefs/RB10073.html (accessed on 18 May 2020).
9. Reding, D.F.; Eaton, J. Trends 2020–2040. Exploring the S&T Edge. *NATO Sci. Technol. Organ.* **2020**. Available online: https://www.nato.int/nato_static_fl2014/assets/pdf/2020/4/pdf/190422-ST_Tech_Trends_Report_2020-2040.pdf (accessed on 9 May 2020).
10. Lye, H. AI, Data, Space and Hypersonics Set to Be ‘Strategic Disruptors’: NATO. *Army Technol.* **2020**. Available online: <https://www.army-technology.com/features/ai-data-space-and-hypersonics-set-to-be-strategic-disruptors-nato/> (accessed on 9 May 2020).
11. Lewin, K. *Field Theory of Social Science. Selected Theoretical Papers*; Cartwright, D., Ed.; Harper & Brothers: New York, NY, USA, 1951.
12. Walden, D.D.; Roedler, G.J.; Forsberg, K.J.; Hamelin, R.D.; Shortell, T.M. *Systems Engineering Handbook. A Guide for System Life Cycle Processes and Activities*, 4th ed.; John Wiley & Sons: New York, NY, USA, 2015.
13. Kelley, H.H. *Personal Relationships: Their Structure and Processes*; Lawrence Earlbaum: Hillsdale, NJ, USA, 1979.
14. Jones, E.E. Major developments in five decades of social psychology. In *Handbook of Social Psychology*; Gilbert, D.T., Fiske, S.T., Lindzey, G., Eds.; McGraw-Hill: Boston, MA, USA, 1998; Volume 1, pp. 3–57.
15. Axelrod, R. *The Evolution of Cooperation*; Basic: New York, NY, USA, 1984.
16. Rand, D.G.; Nowak, M.A. Human cooperation. *Cogn. Sci.* **2013**, *17*, 413–425. [CrossRef]
17. Nosek, B. Open Collaboration of Science: Estimating the reproducibility of psychological science. *Science* **2015**, *349*, 943.
18. Blasch, E.; Cruise, R.; Aved, A.; Majumder, U.; Rovito, T. Methods of AI for Multimodal Sensing and Action for Complex Situations. *AI Mag.* **2019**, *40*, 50–65. [CrossRef]
19. Hare, B.; Woods, V. *Survival of the Friendliest. Understanding our Origins and Rediscovering Our Common Humanity*; Penguin Random House: New York, NY, USA, 2020.
20. Kenny, D.A.; Kashy, D.A.; Bolger, N. Data Analyses in Social Psychology. In *Handbook of Social Psychology*, 4th ed.; Gilbert, D.T., Fiske, S.T., Lindzey, G., Eds.; McGraw-Hill: Boston, MA, USA, 1998; Volume 1, pp. 233–265.
21. Lawless, W.F. Quantum-Like Interdependence Theory Advances Autonomous Human–Machine Teams. *Entropy* **2020**, *22*, 1227. [CrossRef] [PubMed]
22. Berscheid, E.; Reis, H.T. Attraction and close relationships. In *The Handbook of Social Psychology*, 4th ed.; Gilbert, D.T., Lindzey, G., Fiske, S.T., Eds.; Lawrence Erlbaum: London, UK, 1998; pp. 193–281.
23. Lawless, W.F. The Interdependence of Autonomous Human-Machine Teams: The Entropy of Teams, But Not Individuals, Advances Science. *Entropy* **2019**, *21*, 1195. [CrossRef]
24. Bohr, N. Science and the unity of knowledge. In *The Unity of Knowledge*; Leary, L., Ed.; Doubleday: New York, NY, USA, 1955; pp. 44–62.

25. Committee on the Science of Team Science; Board on Behavioral, Cognitive, and Sensory Sciences; Division of Behavioral and Social Sciences and Education; National Research Council. *Enhancing the Effectiveness of Team Science*; Cooke, N.J., Hilton, M.L., Eds.; National Academies Press: Washington, DC, USA, 2015.
26. Cummings, J. *Team Science Successes and Challenges*. National Science Foundation Sponsored Workshop on Fundamentals of Team Science and the Science of Team Science; Bethesda: Rockville, MD, USA, 2015.
27. Hansen, L.P. *How Quantitative Models Can Help Policy Makers Respond to COVID-19*. Good Policy Making Recognizes and Adapts to the Uncertainty Inherent in Model Building; Chicago Booth Review: Chicago, IL, USA, 2020.
28. Markowitz, H.M. Portfolio Selection. *J. Financ.* **1952**, *7*, 77–91.
29. Artzner, P.; Delbaen, F.; Eber, J.M.; Heath, D. Coherent Measures of Risk. *Math. Financ.* **2001**. [CrossRef]
30. Hughes, W. Proposed Merger Holds Promise for U.S. National Security. *Real Clear Def.* **2019**. Available online: https://www.realcleardefense.com/articles/2019/08/28/proposed_merger_holds_promise_for_us_national_security_114704.html (accessed on 28 August 2019).
31. Bisbey, T.M.; Reyes, D.L.; Traylor, A.M.; Salas, E. Teams of psychologists helping teams: The evolution of the science of team training. *Am. Psychol.* **2019**, *74*, 278–289. [CrossRef]
32. Wissner-Gross, A.D.; Freer, C.E. Causal Entropic Forces. *Phys. Rev. Lett.* **2013**, *110*, 1–5. [CrossRef]
33. Tetlock, P.E.; Gardner, D. Superforecasting: The Art and Science of Prediction. *Risks* **2016**, *4*, 24. [CrossRef]
34. Blanton, H.; Jaccard, J.; Klick, J.; Mellers, B.; Mitchell, G.; Tetlock, P.E. Strong Claims and Weak Evidence: Reassessing the Predictive Validity of the IAT. *J. Appl. Psychol.* **2009**, *94*, 567–582. [CrossRef]
35. Shortell, T.; Lawless, W.F. *Case Study. Uber Fatal Accident 2018*. *Systems Engineering Handbook*; International Council on System Engineering (NCOSE): San Diego, CA, USA, 2021. (in preparation)
36. McChrystal, S.A.; Collins, T.; Silverman, D.; Fu, C. *Team of Teams: New Rules of Engagement for a Complex World*; Penguin: New York, NY, USA, 2015; pp. 14–19.
37. Sofge, D.; Mittu, R.; Lawless, W.F. AI Bookie Bet: How likely is it that an AI-based system will self-authorize taking control from a human operator? *AI Mag.* **2019**, *40*, 79–84. [CrossRef]
38. Cooley, M. Human-centred systems. In *Designing Human-Centred Technology: A Cross-Disciplinary Project in Computer-Aided Manufacturing*; Rosenbrock, H.H., Ed.; Springer: London, UK, 1990.
39. Cooley, M. On human-machine symbiosis. In *Cognition, Communication and Interaction. Transdisciplinary Perspectives on Interactive Technology*; Gill, S.P., Ed.; Springer: London, UK, 2008.
40. Norman, D.A. Human-centered design considered harmful. *Interactions* **2005**, *12*, 14–19. Available online: <https://dl.acm.org/doi/pdf/10.1145/1070960.1070976> (accessed on 28 April 2020). [CrossRef]
41. Kissinger, H.A. How the Enlightenment Ends. Philosophically, Intellectually—In Every Way—Human Society is Unprepared for the Rise of Artificial Intelligence. *The Atlantic*. June 2018. Available online: <https://www.theatlantic.com/magazine/archive/2018/06/henry-kissinger-ai-could-mean-the-end-of-human-history/559124/> (accessed on 15 June 2018).
42. Ready for Weapons with Free Will? *New York Times*. Available online: <https://www.nytimes.com/2019/06/26/opinion/weapons-artificial-intelligence.html> (accessed on 15 July 2019).
43. Atherton, K.D. When Should the Pentagon Update Its Rules on Autonomous Weapons? *Artif. Intell.* **2019**. Available online: <https://www.c4isrnet.com/battlefield-tech/2019/12/12/when-should-the-pentagon-update-its-rules-on-autonomous-weapons/> (accessed on 28 April 2020).
44. Freedberg, S.J.J. Project Convergence: Linking Army Missile Defense, Offense, & Space. The Army wants to do a tech demonstration in the southwestern desert—COVID permitting—Of how the new weapons systems it's developing can share data. *Breaking Defence*, 18 May 2020.
45. Hay, W.A. The Cabinet' Review. George Washington and the creation of an American institution. By Lindsay Chervinsky. Belknap/Harvard: Advise and Dissent. Washington Stalked from the Chamber when the Senate Referred his Questions to a Committee. He Decided to Fashion His Own Advisory Group. *Wall Street Journal*, 15 April 2020.
46. Dallen, R.A. *Major Eisenhower: Decision-Making and Consensus in an Unfamiliar Context*; School of Advanced Military Studies; United States Army Command and General Staff College: Fort Leavenworth, KS, USA, 2015.
47. Eisenhower, D. *The Papers of Dwight David Eisenhower*; Galambos, L., Ed.; Johns Hopkins University Press: Baltimore, ML, USA, 1984; Volume 11, p. 1516.
48. Fouad, H.Y.; Raz, A.K.; Llinas, J.; Lawless, W.F.; Mittu, R. Finding the Path toward Design of Synergistic Human-Centric Complex Systems. *IEEE Syst. Man Cybern.* **2021**. (in preparation).
49. Barabási, A.L. Network Science: Understanding the Internal Organization of Complex Systems. In *2012 AAAI Spring Symposium Series, Stanford, CA, USA, 26–28 March 2012*; AAAI Publications: Palo Alto, CA, USA, 26–28 March 2012.
50. Amadae, S.M. Rational Choice Theory. Available online: <https://www.britannica.com> (accessed on 20 May 2018).
51. Liu, Y.Y.; Barabási, A.L. Control principles of complex systems. *Rev. Mod. Phys.* **2016**, *88*, 1–58. [CrossRef]
52. Thagard, P. Cognitive science. In *The Stanford Encyclopedia of Philosophy*; Zalta, E.N., Ed.; The Metaphysics Research Lab: Stanford, CA, USA, 2019.
53. Hansen, L.P. Uncertainty inside and outside of economic models. *The Nobel Prize*. 8 December 2013. Available online: <https://www.nobelprize.org/prizes/economic-sciences/2013/hansen/lecture/> (accessed on 25 April 2020).

54. Mann, R.P. Collective decision making by rational individuals. *Proc. Natl. Acad. Sci. USA* **2018**, *115*, E10387–E10396. [CrossRef] [PubMed]
55. Pearl, J. Reasoning with Cause and Effect. *AI Mag.* **2002**, *23*, 95–111.
56. Pearl, J.; Mackenzie, D. AI Can't Reason Why. The current data-crunching approach to machine learning misses an essential element of human intelligence. *Wall Street Journal*, 18 May 2018.
57. National Transportation Safety Board. Preliminary Report Highway HWY18MH010. The Information in this Report is Preliminary and will be Supplemented or Corrected During the Course of the Investigation. *National Transportation Safety Board*, 24 May 2018.
58. National Transportation Safety Board. Inadequate Safety Culture Contributed to Uber Automated Test Vehicle Crash—NTSB Calls for Federal Review Process for Automated Vehicle Testing on Public Roads. *National Transportation Safety Board*, 19 November 2019.
59. Wilkes, W. How the World's Biggest Companies Are Fine-Tuning the Robot Revolution. Automation is leading to job growth in certain industries where machines take on repetitive tasks, freeing humans for more creative duties. *Wall Street Journal*, 14 May 2018.
60. Boudette, N.E. Can Elon Musk and Tesla Reinvent the Way Cars Are Made? Scrambling to Turn Out Its First Mass-Market Electric Car, The Automaker Set up Multiple Assembly Lines and Is Changing Production Processes on The Fly. *New York Times*, 30 June 2018.
61. Kottenstette, R. Elon Musk Wasn't Wrong about Automating the Model 3 Assembly Line—He Was Just Ahead of His Time. *TechCrunch*, 5 March 2020.
62. Sherman, L. Tesla Survived Manufacturing Hell—Now Comes the Hard Part. *Forbes*, 20 December 2018.
63. Easterbrook, G. By Edward Niedermeyer (2019), 'Ludicrous. The unvarnished story of Tesla Motors. BenBella'. Book Review: A Revolutionary Old Product. For all its innovations—from battery power to computer-assisted driving—Tesla continues to emphasize its cars' speed, luxury and sex appeal. *Wall Street Journal*, 27 August 2019.
64. LaGrone, S. USS Fitzgerald Returns to Sea After Repairs Caused by Fatal 2017 Collision. *US Naval Institute*, 3 February 2020.
65. National Transportation Safety Board. *Collision between US Navy Destroyer John S McCain and Tanker Alnic MC Singapore Strait, 5 Miles Northeast of Horsburgh Lighthouse August 21, 2017*; Marine Accident Report NTSB/MAR-19/01; PB2019-100970 Notation 58325; National Transportation Safety Board: Washington, DC, USA, 2019.
66. Bradbury, J.A.; Branch, K.M.; Malone, E.L. *An Evaluation of DOE-EM Public Participation Programs (PNNL-14200)*; Pacific Northwest National Laboratory: Richland, WA, USA, 2003.
67. Lawless, W.F.; Akiyoshi, M.; Angjellari-Dajci, F.; Whitton, J. Public consent for the geologic disposal of highly radioactive wastes and spent nuclear fuel. *Int. J. Environ. Stud.* **2014**, *71*, 41–62. [CrossRef]
68. Cary, A. New \$13 Billion Contract Awarded for Hanford Tank Farm Cleanup. *Tri-City Herald*, 14 May 2020.
69. Cooke, N. Effective human-artificial intelligence teaming. In Proceedings of the AAAI-2020 Spring Symposium, Stanford, CA, USA, 23–24 March 2020.
70. White Paper. European Governance (COM (2001) 428 Final; Brussels, 25 July 2001). Brussels, Commission of the European Community (doc/01/10, Brussels, 25 July 2001). Available online: https://ec.europa.eu/commission/presscorner/detail/en/DOC_01_10 (accessed on 1 January 2002).
71. Akiyoshi, M.; Lawless, W.F.; Whitton, J.; Charnley-Parry, I.; Butler, W.N. *Effective Decision Rules for Public Engagement in Radioactive Waste Disposal: Evidence from the United States, the United Kingdom, and Japan*, Nuclear Energy; OECD Nuclear Energy Agency Workshop; Nuclear and Social Science Nexus: Paris, France, 2019.
72. Schrödinger, E. Discussion of Probability Relations between Separated Systems. *Proc. Camb. Philos. Soc.* **1935**, *31*, 555–563. [CrossRef]
73. Von Bertalanffy, L. *General System Theory: Foundations, Development, Applications*; Braziller: New York, NY, USA, 1968.
74. Checkland, P. *Systems Thinking, Systems Practice*; Wiley: New York, NY, USA, 1999.
75. Reed, T. Has Boeing Become A Permanent No. 2 To Airbus? *Forbes*, 9 December 2020.
76. Mou, M. HNA Group's Financial Maneuvers Powered Its Rise—and Caused Its Downfall. The findings of a working group are shedding some light on the unorthodox ways that money was moved around within the conglomerate. *Wall Street Journal*, 8 February 2021.
77. Fitch, A. AMD Agrees to Buy Rival Chip Maker Xilinx for \$35 Billion. U.S. Semiconductor Industry Consolidation Gains Pace with Another Landmark Transaction. *Wall Street Journal*, 8 February 2021.
78. Grant, N. Salesforce to Buy Software Maker Slack for \$27.7 Billion. *Bloomberg News*, 1 December 2020.
79. Martyushev, L.M. Entropy and entropy production: Old misconceptions and new breakthroughs. *Entropy* **2013**, *15*, 1152–1170. [CrossRef]
80. Shelbourne, M. HII Purchases Autonomy Company to Bolster Unmanned Surface Business. *USNI News*, 4 January 2021.
81. Kimura, M. *The Neutral Theory of Molecular Evolution*; Cambridge University Press: Cambridge, UK, 1983.
82. Herrera, S.; Chin, K. Amazon, Berkshire Hathaway, JPMorgan End Health-Care Venture Haven. Company had targeted innovations in primary care, insurance coverage, prescription drug costs. *Wall Street Journal*, 4 January 2021.
83. McDowell, J.J. On the theoretical and empirical status of the matching law and matching theory. *Psychol. Bull.* **2013**, *139*, 1000–1028. [CrossRef] [PubMed]
84. McFadden, J.; Al-Khalili, J. The origins of quantum biology. *Proc. R. Soc. A* **2018**, *474*, e20180674. [CrossRef] [PubMed]

-
85. Lawless, W.F. The entangled nature of interdependence. Bistability, irreproducibility and uncertainty. *J. Math. Psychol.* **2017**, *78*, 51–64. [[CrossRef](#)]
 86. Walch, K. The Future with Level 5 Autonomous Cars. *Forbes*, 20 June 2020.
 87. Lawless, W.F. *Interdependence for Human-Machine Teams, Foundations of Science*; Springer: Berlin/Heidelberg, Germany, 2019. [[CrossRef](#)]
 88. Schrödinger, E. *What is Life? The Physical Aspect of the Living Cell*; Dublin Institute for Advanced Studies at Trinity College: Dublin, Ireland, 1944.
 89. Wei, Y.; Wang, X.; Guan, W.; Nie, L.; Lin, Z.; Chen, B. Neural Multimodal Cooperative Learning toward Micro-Video Understanding. *IEEE Trans. Image Process.* **2020**, *29*, 1–14. [[CrossRef](#)]