



The stochastic nature of machine learning and its implications for high-consequence AI

Thilo Stadelmann^{1,2,3} · Philipp H. Merkt^{4,5} · Kasey Barr⁶

Received: 18 October 2025 / Accepted: 9 February 2026 / Published online: 15 March 2026
© The Author(s) 2026

Abstract

Modern AI systems achieve remarkable performance through fundamentally stochastic processes—machine learning models that function as high-dimensional probability density functions, outputting the most likely predictions given training data. While these systems can match or exceed human performance on average, their methodology produces fundamentally different failure modes than human reasoning, leading to errors that appear nonsensical from a human perspective but are predictable given their probabilistic nature. This has critical implications for high-consequence environments such as military applications where decisions cannot be reversed and may affect lives and material assets definitively. Through detailed analysis of contemporary AI's working mechanisms—particularly how knowledge is acquired through statistical pattern recognition rather than causal reasoning—this paper demonstrates why AI systems inherit biases, cannot distinguish plausibility from factual correctness, and exhibit confident behaviour even when wrong. Written to provide guidance for non-technical stakeholders, specifically but not exclusively in the military domain, it posits that for effective deployment of AI in high-consequence scenarios, processes need to be implemented that make sure all human stakeholders are aware of these facts, develop adequate scepticism of the AI system, and remain actively involved in the decision-making. For military applications specifically, this understanding reveals that effective human-AI collaboration requires more than oversight: it demands co-learning frameworks that maintain meaningful human control through bidirectional information flow, and behavioural and functional awareness on the human side. We give an outlook to decentralized, co-learned AI system tailored to specific teams in dedicated co-learning labs to mitigate power concentration risks while preserving essential human capacities, including moral judgment to exercise mercy.

Keywords Artificial intelligence · Military decision-making · Human-AI collaboration · Co-learning · Meaningful human control · Strategic resilience

✉ Thilo Stadelmann
stdm@zhaw.ch
Philipp H. Merkt
philipp.merkt@hs-fresenius.de
Kasey Barr
kbarr@runi.ac.il

¹ Centre for Artificial Intelligence, Zurich University of Applied Sciences, Winterthur, Switzerland

² Fellow, ECLT European Centre for Living Technology, Venice, Italy

³ Member, Scientific Council, IAEAI Israeli Association for Ethics in Artificial Intelligence, Tel Aviv, Israel

⁴ Research and Education Center for Extraordinary Tactical Situations and Strategic Resilience (18_RECESS), Hochschule Fresenius, Idstein, Germany

⁵ Chair for the Education of Personal and Interpersonal Competencies in Healthcare, Faculty of Health (Department of Human Medicine), Witten/Herdecke University, Witten, Germany

⁶ Rubenstein Center for Constitutional Challenges, Reichman University, Herzliya, Israel

1 Introduction

Modern AI has earned a reputation of yielding results comparable to human level performance for a wide array of tasks [1, 2], e.g., visual recognition [3], text and video comprehension [4], decision support based on heterogeneous data analysis [5], and first steps towards autonomous multi-step acting [6]. Indeed, for many benchmarks, AI results even surpass human performance, in line with many anecdotal examples [7]. At the same time, similarly real experiences exhibit uncanny ‘stupid errors’ of AI systems that do not exhibit common sense, making one question bold claims of ‘understanding’, ‘reasoning’, or, generally, ‘fitness for purpose’ of any practical sort of these models [8–16].

This has important ramifications in high-consequence environments such as certain military applications where decisions cannot be taken back and may affect lives and material assets in a definitive way: how shall human operators deal with such fluctuation in order fulfillment by their AI systems? After all, AI has been suggested (and, in current conflicts that usually speed up innovation and adoption: is used) as an important component in aspects ranging from the military decision-making process (MDMP) to lethal autonomous weapon systems (LAWS). For example, Meerveld et al. [17] express the hope that the use of AI could help in every step of the MDMP with automation and support that mitigates human decision-making biases, overcomes human inadequacy to extract knowledge from high volumes of data, and leads to higher efficiency and quality. They also point out specific challenges, like AI systems themselves being not free of biases, or dangers in providing too much autonomy to AI systems, the latter calling for human-AI collaboration as the standard application scenario. Indeed, specific challenges of human-AI collaboration exist (e.g., ensuring reasonable human agency [18]), bias [19] and other risks associated with AI [20] need careful mitigation, and how AI interacts with our humanity [21] will hopefully be a major theme of future thought. But will AI’s use in high-consequence scenarios like the military (civilian uses are also included, e.g., in safety-critical network operations [22, 23]) be automatically to the advantage of respective organizations once such challenges—those that are applicable for a task at hand—are handled well?

In this paper, we argue from the point of view that a *basic understanding of the foundational working mechanisms of this technology is necessary for everyone involved* to know the ramifications of its inner workings on the task at hand—ramifications that manifest themselves for example in the ‘stupid errors’ indicated above (which are to be expected once the methods are comprehended). The following sections hence will provide this understanding (Sec. 2), derive consequences for military and other high-consequential use

cases (Sec. 3), and formulate recommendations (Sec. 4). These recommendations align with the literature on Meaningful Human Control (MHC), which argues that ‘meaningful’ control is not satisfied by nominal human presence or a formal veto. What matters is whether socio-technical systems remain appropriately responsive to human reasons and support responsibility attribution [24–26]. Rather than relitigating the MHC debate, this paper translates the stochastic properties of contemporary machine learning (ML) into practical design, training, and governance implications that bear directly on when oversight is genuinely meaningful.

2 The nature of AI

Artificial intelligence has been defined as the simulation of intelligent behaviour with a computer [2, 27, 28] (cp. the opening hypothesis in [29]). For this, the field of AI, founded in the 1950 s, does not offer a unified theory or methodology—there is no one way to “build AI” (the phrase itself is misleading), nor any known path towards anything resembling ‘artificial general intelligence’ (AGI). Rather, AI holds a toolbox full of different methods that are each appropriate to simulate one or several specific behaviours (cp. the definitive AI textbook by Russell and Norvig [30]).

2.1 Symbolic AI: logic and reasoning

An important part of the AI toolbox are so-called ‘symbolic’ methods: they manipulate abstract symbols (think: variables as in math which stand for some semantic object) using formal logic to implement rigorous reasoning processes. That is, given a knowledge base of ‘facts’ and ‘rules’, those can be used to infer any logically deducible fact that follows from that knowledge. Respective systems like CYC [31] were particularly strong in the 1980’s and 1990’s, fuelled the ‘expert systems’ hype around AI at that time, and have been (and are) used advantageously in high-consequence scenarios since then [32]. For example, the AI system used for logistics planning during ‘Operation Desert Storm’ has been said to have “paid back all of DARPA’s 30 years of investment in AI in a matter of a few months” [33].

Symbolic methods remain important (e.g., today’s navigation systems calculate their wayfinding based on symbolic AI algorithms like A* [34]). Yet, they generally suffer from the complexity of the real world: there is a gap between what can be perceived from (potentially error-contaminated) measurements and the clean and abstracted world of logical descriptions (that even the “person-century effort” to build CYC could not bridge despite useful niche applications). Hence, the focus of AI research & development shifted to methods that operate below the ‘symbol’

level, directly on data, and are able to adapt to it. There is hope that both methodologies can one day be united, but currently, so-called ‘neuro-symbolic’ AI is still in its infancy [35].

2.2 Subsymbolic AI: statistical machine learning

Since the mid-1990 s, the predominant part of the AI toolbox is ML, used wherever the intended behaviour cannot be described by a set of rules (logic). ML’s most successful methods are essentially function approximation [36]: a mapping is sought from data $\vec{x}$ to some outcome $\vec{y}$, which is to be performed by some function $f(x) = y$ (vector notation is commonly dropped). Inputs are usually high-dimensional numeric representations of real-world data: imagine, for example, x to be the concatenation of all the pixel values (each one integer for a greyscale value or 3 integers to encode colour as red-green-blue) of an image, and y a flag

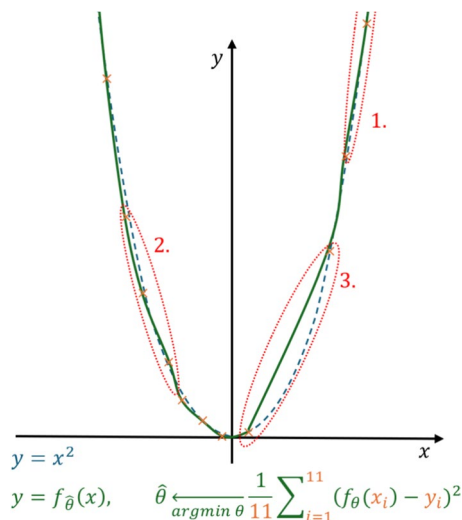


Fig. 1 In blue, dotted: a plot of a parabola $y = x^2$ as a toy example of some real-world event which’s outcomes (y) shall be predicted by AI. A symbolic AI model can be thought of as having access to this true equation, from which it is able to perfectly reproduce each y for any given x . In green, solid: a curve fitted through the 11 orange training data points, resembling the result of training a subsymbolic ML model $y = f_{\hat{\theta}}(x)$ on $\{x_1, x_2, \dots, x_{11}\}$ (counting training instances from left to right) to find optimal parameters $\hat{\theta}$ through minimizing the prediction error for the known training samples. As real-world data, the training points may contain small measurement errors as visible in the figure (i.e., they do not lie perfectly on the blue, dotted parabola); this and other reasons lead to a suboptimal fit (correspondence between the true and the learned model/curve). Consider the red, circled areas: (1) when training data is correct and the used function template for adaptation through training suits the real underlying function, model and true function coincide well. (2) Small measurement errors in the data lead to a suboptimal but likely tolerable fit. (3) In regions with low training data density, no training signal provides guidance for fitting in this area. This leads to larger deviations from the true function; a function with a higher capacity to adapt (like a deep neural network, which can model arbitrarily wiggly functions) could likely zig-zag around wildly between the two far-apart training points x_8 and x_9 .

indicating the presence of some object in the image (1 for “yes”, 0 for “no”) [37]. Or x to be a concatenation of so called ‘word embeddings’ representing a question in natural language, and y being a word embedding for the likely next word (supposedly starting the answer) [38]. Or x consisting of two concatenated structured database entries describing situations (weather, geolocation, other properties; all properly numericized and concatenated), and y being a measure of similarity of the two situations [39].

Evidently, ML is a versatile paradigm: many ‘intelligent behaviours’ can be stated as mapping from an input to some output. Methodically, a human ML engineer provides a suitable ‘function template’ (i.e., a function category that can easily represent the mapping; for example, a straight line cannot represent the dotted curve from Fig. 1, but a polynomial could) as well as a ‘learning algorithm’ that tunes the function template’s parameters to a set of given $\{x, y\}$ pairs called ‘training data’. In recent years, ‘neural networks’, which existed since the field’s inception, rose to unprecedented prominence, becoming the function template of choice for tasks involving perception and cognition [40, 41]. Basically all AI systems that have made the news since 2012 are based on them. This success stems from ‘deep’ neural networks (which consist of several consecutive layers of neurons) that give the function a high capacity to adapt to the training examples: they are general function approximators [42]. Still, the suitability for a given task depends on clever choices of their internal ‘architecture’ and ‘hyperparameters’ (see Segessenmann et al. [21] for an in-depth explanation for non-technical readers).

What principles underly this way of ‘learning’ and are important to understand in order to develop intuition for the nature of ML’s results? ‘Statistical’ learning, as it has been called [43], approximates an unknown, underlying function based on a finite, noisy set of samples. The goal is to interpolate between these given training instances to generalize to novel, previously unseen instances (a process called ‘inductive learning’). Therefore, some parameterizable continuous function is fitted to the training data by systematically adjusting the parameters of that function to minimize a measure of dissimilarity between the predicted and known outcomes ($\hat{y}$ and y , respectively). For neural networks (and many other ML approaches), this optimization process resembles the statistical principle of ‘maximum likelihood’ estimation: the resulting function yields the most *likely* result, given all the evidence present in the training data [44]. For classification tasks (i.e., category prediction), it factually implements $f(x) = p(y|x)$, the conditional probability of the outcome y given the data x . For other tasks like regression (the prediction of continuous numeric values), the model outputs point estimates that represent the most likely values given the training data distribution. In any case, what has been

learned (implicitly or explicitly) is the probability density of outputs given the inputs.

This means that the resulting function (also called ‘model’ in AI and ML) has to be seen as a *probabilistic* function: it predicts a result with a certain likelihood, i.e., involves a measure of uncertainty in the prediction. As Fig. 1 illustrates, this uncertainty might be low in parts of the domain of x with a dense sampling of training data points (and given that (a) the model has been trained on enough data; (b) the chosen function template is suitable for the kind of data and underlying function; and (c) any new instances follow the same underlying distribution as the training data). But it might also be extraordinarily high in areas of the input that are far away from any seen example. As the model implements a continuous mapping, it will still predict a $\hat{y}$, not knowing that it doesn’t know. Also for the developer it is hard to tell in advance how accurate the model will be: generally, ML is an empirical science and there is no way of knowing theoretically how well a specific model will do on a task. Rather, the performance is measured experimentally on a ‘test set’ (a hold-out portion of the original training data), and the result is extrapolated to unseen data under the assumption that these will resemble the training data’s distribution.

A couple of properties of this type of ‘learning’ are notable: first, only the function template’s parameters are ‘learned’ (i.e., fitted to the data); the ‘architecture’ (choice of specific function template), hyperparameters (specific detailed choices in the configuration of the function template and learning setup), and learning algorithm is not part of automatic adaptation. They need to be found by a separate process (typically manual selection by a human, though automation is possible [45]) based on prior knowledge of the problem domain. This knowledge and the algorithmic choices based on them become a necessary part of the model as its ‘inductive bias’—a predetermined idea where and how to look for the patterns the model seeks to pick up (as Mitchell [46] points out, any (also human) learning without this bias is futile). Second, as the model picks up all its knowledge only from the fed training data [47], what is not in the data will not be in the model (e.g., things humans infer using their ‘common sense’), and what was in the data will also be present in the model (e.g., human biases [19], for example through biased judgments present in the $\{x, y\}$ pairs). Third, a ML model usually does not learn continually: while there is a subfield of ML called ‘continual learning’, there are still numerous open challenges to solve [48, 49]. Hence, the vast majority of deployed ML-based systems, for example all widely known GenAI systems like large language models (LLMs), do not learn continually [50] (actually, we are not aware of any that does). Rather, they are iteratively trained on the training data until the model’s fit

is sufficiently good. Then, the parameters are fixed and the model is deployed on its task without any further learning: training and ‘inference’ are completely disjunct phases in the ML life cycle.

2.3 Artificial vs. human intelligence: different means, different errors

From the nature of ML outlined above, it becomes evident why models based on neural networks are currently AI’s best attempt to deal with the uncertainties and the messiness of real-world data. This is true for image and video analysis, text analysis and generation, geospatial data analysis, analysis of satellite and other sensor data, etc. If such a model is trained well enough to find acceptance into any application, it likely works very well on average and for typical inputs. At the same time, because of the statistical nature of the model (that has not learned about truth and facts, but statistical plausibility), a result might be wrong in any given case.

Various approaches exist to quantify and manage this uncertainty, including Bayesian neural networks [51], ensemble methods [52], and calibrated confidence scoring [53]. Active learning frameworks can identify when models encounter unfamiliar inputs [54], while human-in-the-loop systems maintain human oversight at critical decision points [55]. However, these techniques often require significant computational overhead, specialized expertise to implement correctly, are often not part of commercial / existing systems, and still cannot eliminate the fundamental issue: ML models remain probabilistic approximators that can fail confidently in unexpected ways.

To grasp the impact of the fundamental likelihood to err, consider the following example of a ML model for visual inspection [56]: having a reasonable accuracy of, say, 95% on a per-image basis, the use case may involve inspecting larger items that are fed subsequently as individual image patches into the classifier—sometimes up to 30 patches. This makes the performance of the model on a per-item basis look rather underwhelming: the potentially acceptable 5% chance of being wrong per patch (image) accumulates to a $1 - (0.95^{30}) = 78.5\%$ chance of misclassification per item. Put differently: it is to be expected that *every* use of that AI system for visual inspection makes a wrong overall prediction.

To gauge the ramifications of the statistical nature of predictions further, consider the following example of using a LLM [57]: a so-called ‘reasoning’ model has been asked the question “The surgeon, who is the boy’s father, says ‘I can’t operate on this boy, he is my son!’ Who is the surgeon to the boy?” The answer is straight-forward from the question’s text, yet the model replies “The surgeon is the boy’s mother,” which is obviously wrong. But to the model, this

makes actually sense, as it goes on to tell: “the riddle plays on the assumption that a surgeon is male.” Indeed, variations of the question exist abundantly on the web as tests for our own human biases, typically associating males with the role of a surgeon. The model has seen all these during its training (LLMs are trained on almost all text openly accessible on the internet) and learned the utter statistical implausibility of answering anything male to a question that looks remotely like the one above. Consequently, the model gives a plainly wrong answer—but one that is totally *plausible* for any AI system built according to the principles of contemporary ML (which are the best ones we currently have; other forms of ML are conceivable, but not yet mature [58]).

This makes it evident that AI (using any of its methods, including ML) works decidedly different than human intelligence (as is already implied by the definition above, stating that intelligent behaviour is simulated rather than intelligence implemented). While on average possibly better than the mean of human outcomes given a specific task, from the different modes of operation under the hood follows that the remaining errors will also be different: AI systems will commit different errors and exhibit different failure patterns than humans. For example, while humans are ill-equipped to sift through high volumes of heterogeneous data because of sheer information overload, AI systems will also overlook and misinterpret things because of their suboptimal (statistical, not causal/common-sensical) understanding of the world.

The different nature of artificial and human intelligence can be finally illustrated with an analogy of a musician and a DJ: while a DJ simulates certain aspects of creating music very well, their method of music creation through remixing and replaying musicians’ original recordings by design is not general. There are many aspects of music beyond the method of turntables and remixing, e.g., certain genres, playing techniques, and settings for musical performances. For example, a DJ cannot produce what cannot be reached by mixing existing recordings, hence new musical genres like the Grunge of the 1980 s or New Metal in the 1990 s would never emanate from them. Similarly, AI does not simulate the way intelligent human behaviour is produced, but certain carefully designed aspects of human behavioural outcomes, with a very specific method of cleverly interpolating between pre-recorded behaviour samples. This makes respective models good at certain things (for which they have been designed and tested) and bad for most others.

3 Discussion

Summarily, almost all relevant contemporary AI systems are based on ML models that are high dimensional *probability density functions*, which output the most likely predictions

given the input data, leading to likely errors that have decidedly different error patterns than human experts. We will discuss ensuing general implications for any high-consequence use of AI in Sec. 3.1, from which we derive the need for and foundations of meaningful human-AI collaboration in Sec. 3.2–3.3 before outlining the necessity for strategic resilience with respect to novel security issues in Sec. 3.4.

3.1 General implications

While the field of AI has developed various mitigation strategies for issues stemming from the discrepancy between being a probability density function and being perceived as humanlike as outlined above, these approaches address symptoms rather than the underlying statistical nature of ML. This has important ramifications for any operator (and their organization) relying on respective results (predictions):

AI results are not ‘neutral’ Models have picked up human biases via the training data and are ignorant of anything not represented in the data or not representable or inferable by the chosen model.

AI results tend to regurgitate the past Applied in the straight-forward way, they reproduce the most likely pattern found in the training data. This can lead to an impoverishment of strategic decisions [59] and have effects worth of consideration when competing (or conflicting) parties rely on basically similar AI decision support (think of the same advisor working for all parties).

AI results have a certain likelihood of failure; being error-free is not part of the methodology. That predictions are statistically plausible does not prevent them from still possibly being wrong.

It is not known to a model if its current output is right or wrong, and it is difficult to predict the correctness technically. Anyway, any result will be reported with optimistic confidence by an AI system. Results must hence be verified by a human capable of doing so independently.

Human errors and AI errors are very different such that AI systems’ errors might seem very stupid (and hence unexpected) from a human point of view. This stems from the completely different mechanisms by which these results are achieved, even when based on the same data.

Models have picked up human biases via the training data and are ignorant of anything not represented in the data or not representable or inferable by the chosen model.

On the systemic level, it is noteworthy that AI systems are powerful tools wielded (ultimately) by individual humans. This leads to higher concentrations of power in these individuals. In high consequence settings that are characterized by stressfulness and life and death decisions (e.g., military use on the battlefield), misuse of such power

must be prevented. Although this is not new with respect to military staff, AI systems shift the distribution of power in unexpected ways. For example, significant power could fall into the hands of software vendors and model providers (through dependencies) or training data engineers (through changing model behaviour by biasing/poisoning training data), etc.

The next sections will indicate directions for mitigation regarding these implications with increasing focus on uses of AI in military contexts.

3.2 The case for ML literacy, co-learning, and decentralized systems

AI results, which are typically meant to make analyses more precise, e.g. in a military setting with respect to intelligence (in the sense of ‘knowledge gathering’) and targeting [60], are attained differently from human precision and error-prone as pointed out above. Any human stakeholder must be firmly aware of this fact and the underlying reasons to develop healthy ‘scepticism’ regarding AI’s predictions and recommendations. Here, our assertion that a ‘basic understanding of the foundational working mechanisms of this technology is necessary for everyone involved’ (cp. Sec. 1 and Tigard [61]) plays out: with it, one *expects* the above implications in working with such a system.

For example, one is not surprised that a high-confidence recommendation by a purely ML-based *situational awareness* system can turn out sub-optimal, because it fuses information channels in a shallower way a human would [62], namely based on low-level statistical signals rather than underlying meaning or causation. An operator with said basic understanding of ML would hence often check some of those signals that a human typically looks for and compare these results with the system’s to gain further practical understanding for the kind of situations where judgment coincides. The basic understanding thus provides a generalizable layer of initial, realistic expectations that can further be adapted to specific cases through additional study or lived experience. It constrains and guides this adaptation and thus does much more than accelerating the process of arriving at realistic expectations—it ensures them, by changing the human’s internal estimates of what are likely ML outputs based on an understanding of their causes. The following hypothetical examples illustrate this further:

Target prioritization under pattern saturation: An AI system designed to assist with target prioritization ranks objects or locations based on patterns learned from prior conflicts. The system may assign high confidence to a particular site because its features closely match historical training examples associated with hostile activity. From the model’s perspective, the recommendation is statistically

well supported. A human operator familiar with recent changes on the ground may recognize that the same pattern now reflects a civilian logistics function that emerged after the training data was collected, and chose to revisit the ML system’s assessment of it, knowing that the system’s output may not be erroneous in a technical sense: it correctly identifies a familiar pattern. Yet it is operationally misleading because statistical similarity is mistaken for present relevance. This illustrates how ML systems can produce confident but contextually incorrect recommendations, and why oversight requires judgment (based on understanding) rather than simple verification.

Anomaly detection and quiet failure: In anomaly-detection applications, an AI system may flag only what deviates sharply from learned norms. Gradual changes, such as slow shifts in movement patterns or communication behavior, may remain unremarkable to the system while raising concern for experienced analysts. The system does not miss the signal. It never learned to treat slow drift as meaningful. Here, the limitation arises not from lack of data, but from the statistical framing of relevance itself. This type of quiet failure is particularly challenging for both automated detection and human oversight. It can be anticipated, however, if analysts using the AI system understand its working mechanism.

Hence, the typical mode of operation in the abovementioned and other high-consequence application scenarios is to build human-AI ‘teams’,¹ with the final responsibility with the human (in terms of Davidovic [68]: for the purpose of safety, with meaningful human control and judgment as the type of engagement, analysing the whole process). But scepticism (or human oversight) alone is not enough: psychological research has shown [18] that humans need to have meaningful agency in any collaboration, otherwise they cannot help but become bored, reverting to mere mechanical approval without exercising supervision. A remedy is offered by the concept of *co-learning* currently being developed in a European research project² for human-AI collaboration in the high-consequence scenario of operating critical network infrastructures [23]. Co-learning maintains a setup in which with every interaction both the human and the

¹ As with many of the terms used to describe AI systems (e.g., ‘intelligence’, ‘learning’, ‘thinking’), the line for undue antropomorphisation is crossed with ‘teaming’ [63, 64, 65]. We adopt it here as an established technical term when we connect to the relevant literature or want to emphasize that the AI tool in this collaboration acts as a very ‘active appliance’ [66]; otherwise, we prefer the umbrella term ‘human-AI collaboration’: Although duely criticized for the very same reasons [67], broadly accepted better alternatives are lacking beyond ‘human-AI interaction’, which does not convey the ‘active appliance’ aspect and is thus deemed insufficient here: you can also interact, for example, with a hammer, but we are concerned with a different kind of—a different quality or level of—interaction here.

² <https://ai4realnet.eu/>.

machine learn from each other via bidirectional information flow: Not only do the humans provide training feedback to a (continually learning [69]) ML system, but the AI system at the same time provides explainable insights to the human [70] that help them understand and scrutinize decisions better.³ This happens within a long-term, iterative process of co-adaptation through interaction that leads to co-learning. The support of such human learning and active involvement in the decision-making process keeps the human interested, engaged, and maintains their sense of agency.

Also the power concentration issues identified above raise fundamental questions about military AI architecture. Rather than prescriptive solutions, we offer a speculative framework that illustrates how the principles of co-learning and decentralized systems might address these challenges—questions that merit serious research attention in future work (cp. Sec. 4): Consider combat situations where individual combatants may be augmented by AI systems that provide extended situational awareness (through perception based on additional sensors) and recommendations (based on fast and comprehensive data analysis). Here, power issues become important: Centrally controlled systems would be prone to overriding the individual's meaningful human agency and could lead to a remote-controlled human army not too different from a robotic one. A potential—speculative—solution might be the following one: Every (group of) combatants receives their own individual, decentralized AI assistant [74], able to co-learn (cp. [75]).

3.2.1 Illustrative examples for co-learning systems in military practice

How could such a scenario play out, e.g., on the level of a *fireteam* and mitigate some of the ramifications highlighted above? First, a setup would be chosen in which each individual AI system must not be overridable by a central unit (ensuring compliance with the chain of command could be achieved by subjecting it directly to the human team leader). Second, each individual AI system would be fine-tuned to its team by being trained together in exercise and real scenarios, so that the resulting human-AI unit would 'know' and complement each other's *specific* weaknesses (because it has co-learned and thus co-adapted to each other). This makes this AI system, without implying any anthropomorphization,⁴ of personal value to the human

team members and worthless for other combatants (e.g., hostile forces). Thus, heightened risks of power misuse are met with checks and balances through what in a human-human collaboration would be called a joint 'team spirit': For example, AI recommendations on ethics would be more likely to be followed by a human team if a consequence of not complying could be to lose the digital comrade (that might chose to disintegrate if ignored too often). This way, common human coping mechanisms with stress and differing opinions by social means would translate to the AI part of the collaboration.

A second scenario shall serve to illustrate how the postulated co-adaptation benefits the involved humans and the overall quality of the result: the '*Recover*' task within a *Personnel Recovery (PR) operation* (cp. Holewijn [77], especially Fig. 3 there). It comprises the complex and complicated task of identifying with high confidence, on the ground under immense time pressure and adverse circumstances, an isolated member of one's own forces, e.g., within a Combat Search & Rescue mission: The decision is to be taken whether the target is taken up (because the person can be identified with certainty, and the situation is reasonably safe) or left behind (because there are doubts on the person's identity or whether the situation constitutes an ambush). Whether the decision is ultimately taken by the commander on the ground or in the Tactical Operation Center, AI systems can contribute important cues with respect to situational awareness (assessing safety, cp. U.S. Department of Defense [78]) and biometrics (adding to identity recognition of persons that may be wounded or disfigured) by means of their pattern recognition powers, i.e., the AI system's recommendation reflects statistical aggregation across prior examples. But the human may hesitate because of causal and situational reasoning about how the present case departs from those examples. The difference is not one of speed or accuracy alone, but of mode of judgment, and the resulting error risks are different in kind.

How can this judgment be honed and the AI system's cues (that could be incorrect at any time, either due to error or because certain input modalities were not taken into account, like smell) be duely incorporated? We argue: By joint experience gained in *co-learning labs* (see below and

³ Explainability through AI systems is a philosophically hotly debated topic: do technical artifacts even possess what it takes to explain [71, 72]? Certainly they do not, in the human sense, and yet the term is helpful as an umbrella for methods such as Grad-CAM for visual inputs [73], which help humans gain intuition into the 'why' behind respective results, helping with their interpretation.

⁴ As Sec. 2 showed, AI systems 'know' (etc.) in a very different sense than humans do. Why then this anthropomorphism? Apart from hype

and fallacy [76], finding helpful language in the context of AI, which is deliberately designed to mimic human behaviour at the surface, is difficult: it necessitates balancing the painting of useful word pictures for naturally human behavioural contexts while not making the technology appear more than it is. Within the scientific fields of AI/ML, the shared implicit understanding is that most of these terms are used as technical terms with figurative meaning. This leads to problems in interdisciplinary or public dialogue where this practice is not known or shared. Therefore, we caution readers to understand terms figuratively that are commonly used in a human context but appear here in relation to AI systems.

Sec. 3.3.1) will the AI system have (machine-) learned, for example, what specific foci of attention ‘Sgt. Snuffy’, who leads the operation on the ground, has. Hence, the system will use this training to provide her with a tailored prioritization of inputs targeted at complementing this specific human’s judgment in the most efficient way that does not lead to cognitive overload. Through the same joint training, Sgt. Snuffy has learned in what kind of situations the AI recommendations are reliable and what environmental cues are indicative of the system missing something or appearing too confident. As decisions need to be taken in split seconds, this fine tuning of tool and human to their specific cooperation can lead to vital advantages while not making such an AI system (as a technical equipment item) less controllable or safe for the organization deploying it. How did the fine tuning take place? Sgt. Snuffy and her team exercised several training PR missions in a co-learning lab (think of a Urban Warfare Training Center with additional focus on AI), having the respective AI system as part of their equipment. Using it in training gave rise to the mentioned experience for the humans, while the lab environment captured multi-sensory data (e.g., from fixed and helmet cameras, radio, tracking systems installed in the facility, etc.) that have later been used to perform ‘fine tuning’ on the ML models underlying the AI system (cp. [79] and how ML models are fine-tuned for a different purpose by Ruiz and Sell [80]).

3.3 Behavioral and functional foundations of meaningful human control

Building on the co-learning framework discussed above, sustaining MHC in high-consequence environments through co-learning requires that human operators cultivate awareness of both *behavioral* and *functional* factors. This emphasis is aligned with the MHC literature, which argues that ‘meaningful’ control depends on whether the overall human–AI arrangement remains appropriately responsive to human reasons and supports responsibility attribution in practice [24, 25, 81, 82].

Behavioral awareness involves understanding how stress, time pressure, and cognitive bias influence human judgment when interacting with probabilistic systems. Decades of research in behavioral decision science demonstrate that even experts, individually and in groups [83, 84], are prone to overconfidence, anchoring, and inconsistent evaluations under uncertainty [85–87]. High workload and time pressure further degrade attention and vigilance, changing how operators notice, interpret, and respond to automated cues [88]. These conditions also intensify well-documented patterns of automation misuse, including complacency and automation bias, in which users defer to an aid even when it is wrong or context-mismatched [89–91]. Because reliance

is not static, behavioral awareness must include trust calibration: The capacity to maintain appropriately proportional reliance as performance, context, and incentives shift [92].

Functional awareness complements this behavioral awareness: Operators must grasp how machine-learning systems represent uncertainty, how their reliability changes with context, and where their confidence diverges from causal truth [93, 94]. Transparency and explainability are critical to functional awareness as they enable operators to interpret how AI systems reason under uncertainty and challenge their recommendations appropriately [95].

In high-consequence domains, ‘accuracy in testing’ is not sufficient as a proxy for trustworthiness in deployment: Dataset shift, underspecification, and out-of-distribution conditions can produce brittle or unstable behavior that is not visible in routine validation [96–98]. Functional awareness therefore includes (a) understanding calibration limits and what probabilistic confidence does (and does not) mean [99], and (b) understanding why interpretability and explanation are not merely transparency virtues but practical tools for contestation, error detection, and bounded reliance [100, 101]. When operators appreciate both their own cognitive dynamics *and* the variability and boundary conditions of AI performance, human–AI collaboration can evolve from passive oversight into an adaptive process of mutual calibration.

This functional awareness is central to this article’s core claim: Because contemporary ML is stochastic and context-sensitive, meaningful human oversight depends on operators understanding how uncertainty is represented. When reliability shifts across contexts, and where system confidence can be overrepresented, reliance must be a matter of judgment based on clear understanding of the boundaries within which the system can perform with accuracy and predictability. How can this be implemented?

3.3.1 Towards implementing MHC via co-learning

Co-learning scenarios should aim to build *reciprocal situational awareness*: A two-way understanding in which humans learn how AI performance varies across operational conditions, and AI systems (through design, training, interface, and feedback loops) are shaped to anticipate predictable human vulnerabilities and restrictions under stress, ambiguity, and tempo. This concept builds on established work on situation awareness and human–automation interaction, while emphasizing co-adaptation rather than one-directional ‘user training’ alone [88–90].

Returning to the concept of co-learning labs from Sec. 3.2.1, they can now more generally be understood as real-world training and evaluation environments in which operational situations (including high-tempo decision

contexts) are simulated through enactment, enabling teams to practice with AI-enabled decision support while generating structured evidence about where human and machine reliabilities intersect, diverge, or degrade: As noted, machines excel at rapid pattern detection and probabilistic reasoning, yet lack causal comprehension and moral sensitivity. Humans bring contextual judgment and ethical evaluation but are vulnerable to fatigue, stress, framing effects, and diffusion of responsibility. Hence, training environments that expose where human and machine reliability cross and interact need to be established that can foster the kind of co-adaptive awareness needed for safe deployment.

These environments are not simply “training ranges.” They are governance-relevant infrastructures that (a) reveal predictable failure modes (behavioral and functional), (b) support the refinement of interfaces, escalation protocols, and contestation pathways, and (c) produce auditable learning artifacts that can feed test-and-evaluation, doctrine, and accountability. In this sense, co-learning supports MHC by operationalizing the conditions under which oversight is *actually* meaningful: it strengthens the operator’s capacity to contest, recalibrate, and redirect reliance, and it strengthens the system’s design alignment with human reasons, responsibility practices, and institutional review.

The co-learning framework advocated for here is not intended to replace established concepts such as calibrated trust [92]. Rather, the two operate at different analytical levels. Calibrated trust refers to a state of appropriate proportional reliance on an automated system at a given time, based on perceived competence, predictability, and context. Co-learning, by contrast, describes an iterative process through which such calibration is developed and maintained over time. By repeatedly exposing human operators to system behavior across varying conditions—including uncertainty, degradation, and failure, co-learning environments support the ongoing adjustment of reliance as system performance and operational contexts evolve. In this sense, co-learning functions as a practical mechanism for sustaining calibrated trust in settings where static training or one-time validation is insufficient.

The behavioral–functional approach outlined above clarifies why MHC cannot be secured by procedural oversight alone. It requires disciplined operator judgment *and* system-level literacy about how probabilistic models behave at and beyond their domain limits. Then, co-learning labs provide a practical bridge between these requirements by turning abstract commitments to control into trained competencies, validated boundaries, and institutionally usable evidence of when and why reliance is warranted: research on human-autonomy teaming highlights that robust collaboration depends on dynamic trust, transparency, and shared (or, in the case of AI: appropriately aligned) mental models that

allow both sides to anticipate one another’s limitations [93, 102]. Co-learning can be the mechanism for strengthening the conditions of MHC (i.e., sustained engagement, calibrated reliance, and contestability) given well-documented risks of complacency and automation bias in human–automation interaction [89, 90, 92].

Humans remain the central decision-makers, exercising authority most effectively when they understand their own cognitive limits and the probabilistic nature of AI reasoning. In this sense, co-learning can support calibrated agency, enabling mutual adaptation while preserving the moral and operational accountability of human operators. Training environments that reveal where human and machine reliability intersect provide a plausible path to cultivating the co-adaptive awareness required for safe deployment. A behavioral-functional design provides a foundation for responsible co-learning and for ensuring that human judgment remains active, informed, and accountable even within stochastic, high-tempo decision systems.

3.3.2 The limits of human oversight

Human oversight is often treated as an inherent safeguard, yet the human-automation literature shows that oversight can be protective or counterproductive depending on operational and cognitive conditions. Oversight is most likely to *help* when operators can independently evaluate the system’s output, when the task tempo permits verification, and when interfaces support situation awareness and informed contestation rather than passive acceptance [88, 92]. In these conditions, human judgment functions as a meaningful check on probabilistic outputs, particularly under uncertainty or domain shift.

However, oversight can *harm* when cognitive workload, time pressure, or organizational incentives push operators toward shallow review, producing “rubber-stamping” rather than genuine evaluation. Under high perceived automation reliability, users tend to drift into complacency and automation bias, deferring to system outputs even when they are incorrect or context-mismatched [89–91]. This risk is compounded when operators lack functional awareness about what model confidence signals mean, or when the human role becomes supervisory monitoring in low-engagement conditions that predict vigilance decrements [88, 89]. In short, ‘human in the loop’ or ‘human on the loop’ is not a sufficient condition for meaningful control if the loop is cognitively thin or institutionally pressured.

These findings imply that the value of oversight is conditional and predictable. Factors that tend to improve oversight include time and workload margins, clear contestation pathways, feedback that enables calibration, and interface design that supports active verification [92].

3.4 Strategic resilience to mitigate hybrid security threats

ML's proneness to bias, confusing of plausibility with factuality, and signalling a high degree of self-confidence even when producing incorrect results also open up new possibilities for attack and disruption. These become strategically relevant for military use in light of hybrid threats [103]: For example, AI systems intended as aids in the MDMP can be deliberately thrown off balance by manipulating data streams. The hybridization of conflict, in which conventional operations are interwoven with cyberattacks, disinformation, and bioterrorist scenarios, makes AI a double-edged sword in military use: It is both an efficiency enhancer and a target for attack. For effective preventive threat prevention, Badalić [104] hence emphasizes that threats must be addressed *ex ante* before they become effective in combat, which can in part be addressed by meaningful human oversight.

Jonsson and Käihkö [105] expand the view of conflict arenas beyond the battlefield with their approach to non-military warfare. The challenges are particularly acute where open-source intelligence (OSINT) data is used for training and mission evaluation by AI: Open sources are heterogeneous, manipulable, and often express aggressive narratives. When AI systems that operate on the principle of statistical plausibility learn from OSINT data, this can distort the situation assessment in lieu of supposed operational effectiveness (cf. [106]). Again, meaningful human oversight is the required counter measure.

Marquis [107] shows that asymmetric actors are particularly successful in exploiting the weaknesses of superior systems. Freudenberg [108]'s theory of irregular warfare makes it clear that asymmetric warfare operates through methods such as deception, infiltration, or overstretching of enemy forces. Applied to the above analysis of AI error modes, this means that irregular opponents can deliberately provoke the inherent susceptibility to bias of probabilistic models—for example, through 'poisoned' OSINT data or simulated patterns—in order to gain disproportionate influence over operational decisions.

To counter these threats systemically, strategic resilience is required, which can be located on three levels [109]: at the micro level among soldiers, analysts, and AI assistance systems; at the meso level among military organizations and networks; and at the macro level among the state, society, and the international order. Resilience means more than just resistance: it encompasses the ability to address threats *ex ante* preventively, *ex nunc* at the moment of action, and *ex post* in the sense of organizational learning. Constellation analysis [110] is a suitable methodological tool for this purpose, as it reveals the interactions between actors, means,

goals, and dynamics and, in conjunction with the hybrid methods of asymmetric warfare described by Freudenberg, shows where AI systems are vulnerable and how their integration into military constellations can be designed responsibly (cf. [111]).

4 Conclusions

The non-negligible likelihood of AI errors in any one situation necessitates the implementation of processes to ensure human operators participating in human-AI collaboration understand the failure modes of their tools and are properly integrated in decision-making. Co-learning schemes can help herewith and at the same time train respective AI systems to compensate for specific errors and limitations of 'their' human users. Furthermore, AI is not only a tool, but also a potential target for attacks in asymmetric conflicts. Only when meaningful human oversight and judgment, decentralized architectures, and resilient organizational set-ups work together can probabilistic systems be prevented from becoming targets and catalysts for strategic misjudgments in highly consequential situations.

4.1 Limitations and future work

This article has argued that meaningful oversight based on behavioural and functional (including ML) awareness is necessary for high-consequence AI deployment, but it does not claim that oversight is uniformly beneficial across contexts (see Sec. 3.3.2). The effectiveness of oversight is conditional and depends on operational constraints, human cognitive capacity, and task properties. The following limitations also address directions for future work:

Human capacity constraints: Human oversight can degrade under bandwidth limits, cognitive load, and information volume. In high-tempo settings, monitoring roles can predict vigilance decrements and "rubber-stamping," especially when verification is effortful and accountability is diffuse. Thus, future work should model oversight as a scarce resource and specify minimum conditions for substantive review (for example, time-to-verify margins, workload thresholds, and interface support for contestation). At the same time, succeeding to implement co-learning frameworks is expected to mitigate a larger proportion of this issue through the improved bi-directional information flow between human and AI, mediated by mutual adaptation.

Expertise differentiation: Oversight competence is likely heterogeneous. Domain expertise and what we term *functional awareness* (ability to interpret uncertainty, calibration limits, and domain shift) may shape error detection, but the distribution and magnitude of these effects are not

established here. Future research should distinguish (a) domain expertise (operational knowledge), (b) system literacy (knowing what confidence does and does not mean), and (c) procedural expertise (ability to execute verification and escalation under time pressure), and test which combinations predict reliable contestation of AI outputs.

Expertise transmission: The co-learning laboratories proposed in this article assume that at least some oversight-relevant skills can be trained, but the pedagogical limits are not yet well characterized. Training may improve calibration and contestation, yet it may also induce overconfidence or brittle heuristics. A priority research agenda is therefore to evaluate training regimes experimentally, with outcomes that include calibration, detection of out-of-distribution failure, and appropriate override behavior under stress and time pressure.

Failure modes of oversight: Moreover, human involvement can worsen outcomes when deliberation introduces harmful delays; when cognitive biases dominate under uncertainty; when users over-ride correct AI outputs due to misplaced confidence; or when the human lacks the domain basis to evaluate the recommendation. These risks reinforce a central point of this article: “human-in-the-loop” is not a sufficient condition for meaningful control. Oversight must be designed so that review is feasible, contestation pathways are clear, and the human role remains cognitively substantive rather than formal. Co-learning labs are proposed as one mechanism for identifying these boundary conditions in practice and for generating auditable evidence about where reliance should be bounded, deferred, or escalated.

In addition to these human factors, open *questions on the ML side* need to be addressed thoroughly in the future to realize co-learning enabled MHC in high-consequence scenarios: how to transform *any* practical ML system used in the respective high-consequence context into one that is continuously learning, overcoming memory issues and catastrophic forgetting, and outputting meaningful confidence estimates? How to enable the information flow from AI system to human necessary for human adaptation via XAI, especially for data modalities other than vision? What ML algorithms are appropriate for the machine part of co-learning, and how to evaluate the joint learning progress? Finally: What aspects of psychology need to be incorporated to which degree to facilitate the human learning?

4.2 Recommendations

Because AI processing routinely exceeds human review capacity, meaningful verification must rely on selective validation strategies that target anomalies, boundary conditions, and irreversible decisions, rather than comprehensive replication of machine analysis. The following principles

are offered as practical orientation precisely because oversight is conditional; they are meant to guide governance toward contexts where oversight is feasible and meaningful, and away from procedural ‘human presence’ that does not improve outcomes.

To keep the discussion actionable, the technical realities outlined in this article can be translated into practical high-level guiding principles for responsible innovation, forming the acronym ‘*GUARD*’:

Governance Assign clear decision authority and responsibility across the AI lifecycle, including who may deploy, modify, pause, or override the system, and under what conditions.

Uphold Human Dignity Ensure that high-stakes decisions remain responsive to human reasons and do not reduce persons to data points; where consequences are severe, preserve the ability to pause, reconsider, or defer.

Anticipate Error Treat mistakes and misfit as expected under uncertainty and shifting conditions; build practices that help users recognize when outputs may be unreliable and require additional scrutiny; teach basic ML literacy to all stakeholders.

Retain Human Agency Structure human involvement as active judgment rather than passive monitoring, including meaningful opportunities to question, contest, and redirect reliance on the system; *create, with high priority, co-learning labs to train this.*

Document Accountability Maintain traceable records sufficient to reconstruct how the system’s output was used and how human judgment was exercised, supporting audit, learning, and responsibility attribution.

Table 1 operationalizes *GUARD*. Absent these safeguards, reliance is more likely to drift from proportional, context-sensitive use toward routine and procedural acceptance. These principles provide a practical orientation for embedding AI decision support in ways that keep oversight meaningful, while leaving the operational specification of thresholds and procedures for future research and evaluation.

Not least, designing, developing, and deploying AI systems (e.g., the speculative ones of Sec. 3.2.1) according to these principles of decentralized, co-learned AI systems, specifically in military contexts, would ensure a proper place for the often unwanted but ultimately important human trait of having mercy.

Table 1 Operationalizing GUARD. The oversight requirements resemble the short description above. The implementation conditions give further insight towards operationalization

GUARD principle	Oversight requirement	Implementation conditions
Governance	Clarify decision authority and responsibility across the AI lifecycle (deployment, modification, suspension, override), including conditions under which each applies	Make authority legible: specify who authorizes use, who can pause/suspend, who can override, and who is responsible for review when operating conditions shift or uncertainty rises (cp. [112])
Uphold Human Dignity	Ensure high-stakes decisions remain responsive to human reasons and do not reduce persons to data points; where consequences are severe, preserve the ability to pause, reconsider, or defer	Require a brief human-reasons statement for high-consequence outputs (one sentence explaining the human-relevant grounds beyond the model's output) and preserve a protected pause/deferral pathway where consequences are irreversible or contested
Anticipate Error	Treat mistakes and misfit as expected under uncertainty and shifting conditions; build practices that help users recognize when outputs may be unreliable and require additional scrutiny	Teach basic ML literacy to all stakeholders; couple outputs to uncertainty awareness: require an uncertainty/confidence representation (where available), and include a short “conditions changed?” check (e.g., missing data, degraded inputs, novelty, time pressure) that triggers heightened verification, second review, or deferral
Retain Human Agency	Structure human involvement as active judgment rather than passive monitoring, including meaningful opportunities to question, contest, and redirect reliance on the system	<i>Create, with high priority, co-learning labs to train this</i> ; preserve contestability by design: for designated high-consequence outputs, require an active verification move (cross-check, alternative hypothesis check, or second reviewer) so reliance is not reduced to rubber-stamping under tempo or perceived system authority
Document Accountability	Maintain traceable records sufficient to reconstruct how the system's output was used and how human judgment was exercised in order to support auditing, learning, and responsibility attribution	Maintain a minimal standard of process preservation to include output consulted (and uncertainty representation, if present) and a brief rationale noting salient constraints (time pressure, missing inputs, uncertainty); this can be automated in co-learning lab environments

Acknowledgements We are grateful for the helpful discussions with and input received from Helge Stadelmann, Julia Usher, Ricardo Charriaga, and Manuel Renold.

Author Contributions Conceptualization: TS, KB. Methodology: TS, KB. Investigation: TS, PHM, KB. Writing—original draft preparation: TS. Writing—review and editing: TS, PHM, KB. Project administration: TS.

Funding Aspects of this work have been developed within the AI4REALNET project, which has received funding from European Union's Horizon Europe Research and Innovation programme under the Grant Agreement No 101119527 and from the Swiss State Secretariat for Education, Research and Innovation (SERI). Views and opinions expressed are, however, those of the author(s) only and do not necessarily reflect those of the European Union and SERI. Neither the European Union nor the granting authority can be held responsible for them.

Data availability No datasets were generated or analysed during the current study.

Declarations

Conflict of interest The authors declare no Conflict of interest.

Consent for participation Not applicable.

Informed consent Not applicable.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless

indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

1. Stadelmann, T., Tolkachev, V., Sick, B., Stampfli, J., Dürr, O.: Beyond ImageNet: deep learning in industrial practice. In: Applied data science: lessons learned for the data-driven business, 205–232 (Springer, 2019).
2. Stadelmann, T.: A guide to AI. Global Resilience White Papers (2025)
3. Žigulić, N., Glučina, M., Lorencin, I., Matika, D.: Military decision-making process enhanced by image detection. Information . (2024). <https://doi.org/10.3390/info15010011>
4. Tang, Y., Bi, J., Xu, S., Song, L., Liang, S., Wang, T., Zhang, D., An, J., Lin, J., Zhu, R., Vosoughi, A., Huang, C., Zhang, Z., Liu, P., Feng, M., Zheng, F., Zhang, J., Luo, P., Luo, J., Xu, C.: Video understanding with large language models: A survey. In: IEEE Transactions on Circuits and Systems for Video Technology, pp 1–1. <https://doi.org/10.1109/TCSVT.2025.3566695> (2025)
5. Huang, J., Xu, Y., Wang, Q., Wang, Q.C., Liang, X., Wang, F., Zhang, Z., Wei, W., Zhang, B., Huang, L., et al.: Foundation models and intelligent decision-making: Progress, challenges, and perspectives. The Innovation (2025)
6. Sager, P.J., Meyer, B., Yan, P., von Wartburg-Kottler, R., Etaiwi, L., Enayati, A., Nobel, G., Abdulkadir, A., Grewe, B.F., Stadelmann, T.: A comprehensive survey of agents for computer use: Foundations, challenges, and future directions. Preprint at [arXiv:2501.16150](https://arxiv.org/abs/2501.16150) (2025)
7. Bubeck, S., Chandrasekaran, V., Eldan, R., Gehrke, J., Horvitz, E., Kamar, E., Lee, P., Lee, Y.T., Li, Y., Lundberg, S., et al.:

- Sparks of artificial general intelligence: Early experiments with GPT-4. Preprint at [arXiv:2303.12712](https://arxiv.org/abs/2303.12712) (2023)
8. Brooks, R.: The seven deadly sins of AI predictions. MIT Technology review (2017). <https://www.technologyreview.com/2017/10/06/241837/the-seven-deadly-sins-of-ai-predictions/>
 9. Marcus, G.: Deep learning: A critical appraisal. Preprint at [arXiv:1801.00631](https://arxiv.org/abs/1801.00631) (2018)
 10. von der Malsburg, C., Stadelmann, T., Grewe, B.F.: A theory of natural intelligence. Preprint at [arXiv:2205.00002](https://arxiv.org/abs/2205.00002) (2022)
 11. Neururer, D., Dellwo, V., Stadelmann, T.: Deep neural networks for automatic speaker recognition do not learn supra-segmental temporal features. *Pattern Recogn. Lett.* **181**, 64–69 (2024)
 12. Kambhampati, S.: Can large language models reason and plan? *Ann. N. Y. Acad. Sci.* **1534**(1), 15–18 (2024)
 13. Kambhampati, S., Stechly, K., Valmeekam, K., Saldyt, L., Bhambrri, S., Palod, V., Gundawar, A., Samineni, S.R., Kalwar, D., Biswas, U.: Stop anthropomorphizing intermediate tokens as reasoning/thinking traces! Preprint at [arXiv:2504.09762](https://arxiv.org/abs/2504.09762) (2025)
 14. Narayanan, A., Kapoor, S.: AI as normal technology. 25–09 Knight First Amend. Inst. (2025), <https://perma.cc/HVN8-QGQY> (2025)
 15. Kumar, A., Clune, J., Lehman, J., Stanley, K.O.: Questioning representational optimism in deep learning: The fractured entangled representation hypothesis. Preprint at [arXiv:2505.11581](https://arxiv.org/abs/2505.11581) (2025)
 16. Silver, D., Sutton, R.S.: Welcome to the era of experience. In: *Designing an Intelligence*, MIT Press (2025)
 17. Meerfeld, H.W., Lindelauf, R., Postma, E.O., Postma, M.: The irresponsibility of not using AI in the military. *Ethics Inf. Technol.* **25**(1), 14 (2023)
 18. Waefler, T., Hamouche, S., Eisenegger, A.: The supportive ai framework: from recommending to supporting. In: Schmorow, D.D., Fidopiastis, C.M. (eds.) *Augmented cognition*, pp. 303–317. Springer Nature Switzerland, Cham (2025)
 19. Glüge, S., Amirian, M., Flumini, D., Stadelmann, T.: How (not) to measure bias in face recognition networks. In: *IAPR Workshop on Artificial Neural Networks in Pattern Recognition*, 125–137 (Springer, 2020)
 20. Stademann, T.: Debate: Evidence-based AI risk assessment for public policy. *Public Money & Management* (2025)
 21. Segessenmann, J., Stadelmann, T., Davison, A., Dürr, O.: Assessing deep learning: a work program for the humanities in the age of artificial intelligence. *AI Ethics* **5**(1), 1–32 (2025)
 22. Roost, D., Meier, R., Huschauer, S., Nygren, E., Egli, A., Weiler, A., Stadelmann, T.: Improving sample efficiency and multi-agent communication in RL-based train rescheduling. In: 2020 7th Swiss Conference on Data Science (SDS), IEEE, pp 63–64 (2020)
 23. Mussi, M., Metelli, A.M., Restelli, M., Losapio, G., Bessa, R.J., Boos, D., Borst, C., Leto, G., Castagna, A., Chavarriaga, R., Dias, D., Egli, A., Eisenegger, A., El Manyari, Y., Fuxjäger, A., Galdes, J., Hamouche, S., Hassouna, M., Lemetayer, B., Leyli-Abadi, M., Liessner, R., Lundberg, J., Marot, A., Meddeb, M., Schiaffonati, V., Schneider, M., Stadelmann, T., Usher, J., Van Hoof, H., Viebahn, J., Waefler, T., Zanotti, G.: Human-AI interaction in safety-critical network infrastructures. *iScience* p 113400 (2025). <https://doi.org/10.1016/j.isci.2025.113400>
 24. Santoni de Sio, F., Van den Hoven, J.: Meaningful human control over autonomous systems: a philosophical account. *Front. obotics AI* **5**, 323836 (2018)
 25. Mecacci, G., Santoni de Sio, F.: Meaningful human control as reason-responsiveness: the case of dual-mode vehicles. *Ethics Inf. Technol.* **22**(2), 103–115 (2020)
 26. Veluwenkamp, H.: Reasons for meaningful human control. *Ethics Inf. Technol.* **24**(4), 51 (2022)
 27. Merriam-Webster: Artificial Intelligence. Internet Archive, version April 17 (2021). <https://web.archive.org/web/20210417130256/https://www.merriam-webster.com/dictionary/artificial%20intelligence>
 28. Fuchs, T.: Understanding sophia? On human interaction with artificial agents. *Phenomenol. Cogn. Sci.* **23**(1), 21–42 (2024)
 29. McCarthy, J., Minsky, M.L., Rochester, N., Shannon, C.E.: A proposal for the Dartmouth summer research project on artificial intelligence. Research proposal (1955)
 30. Russell, S.J., Norvig, P.: *Artificial intelligence: a modern approach*, 4th edition. Pearson (2022)
 31. Lenat, D.B.: CYC: A large-scale investment in knowledge infrastructure. *Commun. ACM* **38**(11), 33–38 (1995). <https://doi.org/10.1145/219717.219745>
 32. Nilsson, N.J.: *The quest for artificial intelligence*. (Cambridge University Press, 2009)
 33. Hedberg, S.: DART: Revolutionizing logistics planning. *IEEE Intell. Syst.* **17**(3), 81–83 (2002). <https://doi.org/10.1109/MIS.2002.1005635>
 34. Hart, P.E., Nilsson, N.J., Raphael, B.: A formal basis for the heuristic determination of minimum cost paths. *IEEE Trans. Syst. i. Cybern.* **4**(2), 100–107 (1968)
 35. Bhuyan, B.P., Ramdane-Cherif, A., Tomar, R., Singh, T.: Neuro-symbolic artificial intelligence: a survey. *Neural Comput. Appl.* **36**(21), 12809–12844 (2024)
 36. Jordan, M.I., Mitchell, T.M.: Machine learning: trends, perspectives, and prospects. *Science* **349**(6245), 255–260 (2015). <https://doi.org/10.1126/science.aaa8415>
 37. Kaya, M., Bilge, H.S.: Deep metric learning: A survey. *Symmetry* **11**(9), 1066 (2019)
 38. Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., Sutskever, I., et al.: Language models are unsupervised multitask learners. *OpenAI Blog* **1**(8), 9 (2019)
 39. Schmidhuber, J.: Deep learning in neural networks: an overview. *Neural Netw.* **61**, 85–117 (2015)
 40. LeCun, Y., Bengio, Y., Hinton, G.: Deep learning. *Nature* **521**(7553), 436–444 (2015)
 41. Hornik, K., Stinchcombe, M., White, H.: Multilayer feedforward networks are universal approximators. *Neural Netw.* **2**(5), 359–366 (1989)
 42. Vapnik, V.N.: An overview of statistical learning theory. *IEEE Trans. Neural Netw.* **10**(5), 988–999 (1999)
 43. Prince, S.J.: *Understanding deep learning*. MIT press (2023)
 44. Tuggeren, L., Amirian, M., Rombach, K., Lörrwald, S., Varlet, A., Westernmann, C., Stadelmann, T.: Automated machine learning in practice: state of the art and recent results. In: 2019 6th Swiss Conference on Data Science (SDS), IEEE, 31–36 (2019)
 45. Mitchell, T.M.: *Machine learning*. McGraw-Hill New York (1997)
 46. Stadelmann, T., Klamt, T., Merkt, P.H.: Data centrism and the core of data science as a scientific discipline. *Arch. Data Sci. Series A*. **8**(2) (2022)
 47. Purushwalkam, S., Morgado, P., Gupta, A.: The challenges of continuous self-supervised learning. In: *European conference on computer vision*, 702–721 (Springer, 2022)
 48. Verwimp, E., Aljundi, R., Ben-David, S., Bethge, M., Cossu, A., Gepperth, A., Hayes, T.L., Hüllermeier, E., Kanan, C., Kudithipudi, D., Lampert, C.H., Mundt, M., Pascanu, R., Popescu, A., Tolias, A.S., van de Weijer, J., Liu, B., Lomonaco, V., Tuytelaars, T., van de Ven, G.M.: Continual learning: Applications and the road forward. *Trans. Mach. Learn. Res.* (2024)
 49. Kontogianni, T., Yue, Y., Tang, S., Schindler, K.: Is continual learning ready for real-world challenges? Preprint at [arXiv:2402.10130](https://arxiv.org/abs/2402.10130) (2024)

51. Wang, H., Yeung, D.Y.: A survey on Bayesian deep learning. *ACM Comput. Surv. (csur)* **53**(5), 1–37 (2020)
52. Tuggerer, L., Emberger, R., Ghosh, A., Sager, P., Satyawand, Y.P., Montoya, J., Goldschagg, S., Seibold, F., Gut, U., Ackermann, P., et al.: Real world music object recognition. *Trans. Int. Soc. Music Inf. Retr.* **7**(1), 1–14 (2024)
53. Tian, K., Mitchell, E., Zhou, A., Sharma, A., Rafailov, R., Yao, H., Finn, C., Manning, C.D.: Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback. In: *Proceedings of EMNLP* (2023)
54. Nguyen, V.L., Shaker, M.H., Hüllermeier, E.: How to measure uncertainty in uncertainty sampling for active learning. *Mach. Learn.* **111**(1), 89–122 (2022)
55. Zanzotto, F.M.: Human-in-the-loop artificial intelligence. *J. Artif. Intell. Res.* **64**, 243–252 (2019)
56. Stadelmann, T., Amirian, M., Arabaci, I., Arnold, M., Duives-teijn, G.F., Elezi, I., Geiger, M., Lörrwald, S., Meier, B.B., Rom-bach, K., et al.: Deep learning in the wild. In: *IAPR Workshop on Artificial Neural Networks in Pattern Recognition*, Springer, 17–38 (2018)
57. Stadelmann, T.: How not to fear AI. *TEDxZHAW* (Apr. 10, 2025), <https://stdm.github.io/How-not-to-fear-AI/> (2025)
58. Sager, P.J., Deriu, J.M., Grewe, B.F., Stadelmann, T., von der Malsburg, C.: The cooperative network architecture: Learning structured networks as representation of sensory patterns. *arXiv preprint arXiv:2407.05650* (2025)
59. Stigler, S.M.: Regression towards the mean, historically considered. *Stat. Methods Med. Res.* **6**(2), 103–114 (1997)
60. King, A.: Digital targeting: artificial intelligence, data, and military intelligence. *J. Global Secur. Stud.* **9**(2), oga009 (2024)
61. Tigard, D.W.: On bullshit, large language models, and the need to curb your enthusiasm. *AI and Ethics* pp 1–11 (2025)
62. Liu, M., Wei, J., Liu, Y., Davis, J.: Human and ai perceptual differences in image classification errors. *Proceedings of the AAAI Conference on Artificial Intelligence* **39**(13), 14318–14326 (2025). <https://doi.org/10.1609/aaai.v39i13.33568>
63. Seeber, I., Bittner, E., Briggs, R.O., de Vreede, T., de Vreede, G.J., Elkins, A., Maier, R., Merz, A.B., Oest-Reiß, S., Randrup, N., Schwabe, G., Söllner, M.: Machines as teammates: a research agenda on ai in team collaboration. *Inf. Manag.* **57**(2), 103174 (2020). <https://doi.org/10.1016/j.im.2019.103174>
64. National Academies of Sciences, Engineering, and Medicine.: Human-AI teaming: State-of-the-art and research needs. The National Academies Press, Washington DC. <https://doi.org/10.17226/26355> (2021)
65. Gunkel, D.J., Wales, J.J.: Debate: what is personhood in the age of ai? *AI Soc.* **36**(2), 473–486 (2021)
66. Shneiderman, B.: *Human-centered AI*. Oxford University Press (2022)
67. Evans, K.D., Robbins, S.A., Bryson, J.J.: Do we collaborate with what we design? *Top. Cogn. Sci.* **17**(2), 392–411 (2025). <https://doi.org/10.1111/tops.12682>
68. Davidovic, J.: Rethinking human roles in AI warfare. *Nat. Mach. Intell.* (2025). <https://doi.org/10.1038/s42256-025-01123-6>
69. Wang, L., Zhang, X., Su, H., Zhu, J.: A comprehensive survey of continual learning: theory, method and application. *IEEE Trans. Pattern Anal. Mach. Intell.* **46**(8), 5362–5383 (2024)
70. Dwivedi, R., Dave, D., Naik, H., Singhal, S., Omer, R., Patel, P., Qian, B., Wen, Z., Shah, T., Morgan, G., et al.: Explainable AI (XAI): Core ideas, techniques, and solutions. *ACM Comput. Surv.* **55**(9), 1–33 (2023)
71. O'Hara, K.: Explainable ai and the philosophy and practice of explanation. *Comput. Law Secur. Rev.* **39**, 105474 (2020). <https://doi.org/10.1016/j.clsr.2020.105474>
72. Mattioli, M., Cinà, A.E., Pelillo, M.: Understanding xai through the philosopher's lens: A historical perspective. Preprint at [arXiv:2407.18782](https://arxiv.org/abs/2407.18782) (2024)
73. Selvaraju, R.R., Cogswell, M., Das, A., Vedantam, R., Parikh, D., Batra, D.: Grad-cam: Visual explanations from deep networks via gradient-based localization. In: *Proceedings of the IEEE international conference on computer vision*, 618–626 (2017)
74. Zhu, C., Dastani, M., Wang, S.: A survey of multi-agent deep reinforcement learning with communication. *Auton. Agent. Multi-Agent Syst.* **38**(1), 4 (2024)
75. van den Bosch, K., Schoonderwoerd, T., Blankendaal, R., Neerinx, M.: Six challenges for human-AI co-learning. In: *International Conference on Human-Computer Interaction*, Springer, 572–589 (2019)
76. Placani, A.: Anthropomorphism in AI: hype and fallacy. *AI Ethics* **4**(3), 691–698 (2024)
77. Holewijn, B.: *Personnel recovery - a primer*. Joint Air Power Competence Centre (JAPCC) (2011)
78. US Department of Defense.: *DD Form 1833: Isolated Personnel Report (ISOPREP)*. Department of Defense Form, issued by the U.S. DoD; earliest public form copy dated May 2008 (2008)
79. Koedijk M, Landman A, Bottenheft C, Fonken YM, Binsch O (2026) A virtual reality test to evaluate dismounted soldiers' cognitive and psychomotor performance in an operationally relevant setting. *Frontiers in Psychology* Volume 16 - 2025, DOI 10.3389/fpsyg.2025.1540936
80. Ruiz, D.C., Sell, J.: Fine-tuning and evaluating open-source large language models for the army domain. Preprint at [arXiv:2410.20297](https://arxiv.org/abs/2410.20297) (2024)
81. Chengeta, T.: Defining the emerging notion of meaningful human control in weapon systems. *NYUJ Int'l L & Pol* **49**, 833 (2016)
82. Ekelhof, M.: Moving beyond semantics on autonomous weapons: Meaningful human control in operation. *Global Pol.* **10**(3), 343–348 (2019)
83. Barr, K., Mintz, A.: Groupthink, polythink, and con-div: Identifying group decision-making dynamics. In: *Routledge Handbook of Foreign Policy Analysis Methods*, Routledge, 269–287 (2022)
84. Barr, K., Mintz, A.: Public policy perspective on group decision-making dynamics in foreign policy. *Policy Stud. J.* **46**, S69–S90 (2018)
85. Kahneman, D., Tversky, A.: Prospect theory: an analysis of decision under risk. *Econometrica* **47**(2), 263–291 (1979)
86. Kahneman, D., Sibony, O., Sunstein, C.R. *Noise: A flaw in human judgment*. Hachette UK (2021)
87. Dror, I.E.: Cognitive and human factors in expert decision making: six fallacies and the eight sources of bias. *Anal. Chem.* **92**(12), 7998–8004 (2020)
88. Endsley, M.R.: Toward a theory of situation awareness in dynamic systems. *Hum. Factors* **37**(1), 32–64 (1995). <https://doi.org/10.1518/001872095779049543>
89. Parasuraman, R., Manzey, D.H.: Complacency and bias in human use of automation: an attentional integration. *Hum. Factors* **52**(3), 381–410 (2010)
90. Parasuraman, R., Riley, V.: Humans and automation: use, misuse, disuse, abuse. *Hum. Factors* **39**(2), 230–253 (1997)
91. Merritt, S.M., Ilgen, D.R.: Not all trust is created equal: dispositional and history-based trust in human-automation interactions. *Hum. Factors* **50**(2), 194–210 (2008)
92. Lee, J.D., See, K.A.: Trust in automation: designing for appropriate reliance. *Hum. Factors* **46**(1), 50–80 (2004)
93. Lyons, J.B., Sycara, K., Lewis, M., Capiola, A.: Human-autonomy teaming: definitions, debates, and directions. *Front. Psychol.* **12**, 589585 (2021). <https://doi.org/10.3389/fpsyg.2021.589585>
94. Gao, Q., Xu, W., Shen, M., Gao, Z.: Agent teaming situation awareness (atsa): A situation awareness framework for human-ai teaming. Preprint at [arXiv:2308.16785](https://arxiv.org/abs/2308.16785) (2023)

95. Miller, T.: Explanation in artificial intelligence: insights from the social sciences. *Artif. Intell.* **267**, 1–38 (2019)
96. Ovadia, Y., Fertig, E., Ren, J., Nado, Z., Sculley, D., Nowozin, S., Dillon, J., Lakshminarayanan, B., Snoek, J.: Can you trust your model's uncertainty? evaluating predictive uncertainty under dataset shift. *Adv. Neural Inf. Process. Syst.* **32** (2019)
97. D'Amour, A., Heller, K., Moldovan, D., Adlam, B., Alipanahi, B., Beutel, A., Chen, C., Deaton, J., Eisenstein, J., Hoffman, M.D., et al.: Underspecification presents challenges for credibility in modern machine learning. *J. Mach. Learn. Res.* **23**(226), 1–61 (2022)
98. Hendrycks, D., Gimpel, K.: A baseline for detecting misclassified and out-of-distribution examples in neural networks. In: *International Conference on Learning Representations* (2017)
99. Guo, C., Pleiss, G., Sun, Y., Weinberger, K.Q.: On calibration of modern neural networks. In: *International conference on machine learning*, PMLR, 1321–1330 (2017)
100. Rudin, C.: Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nat. Mach. Intell.* **1**(5), 206–215 (2019)
101. Doshi-Velez, F., Kim, B.: Towards a rigorous science of interpretable machine learning. Preprint at [arXiv:1702.08608](https://arxiv.org/abs/1702.08608) (2017)
102. O'Neill, T., McNeese, N., Barron, A., Schelble, B.: Human-autonomy teaming: a review and analysis of the empirical literature. *Hum. Factors* **64**(5), 904–938 (2022)
103. Kambouris, M.E. *Hybrid Warfare 2.2: Where Biothreats Meet Irregular Operations and Cyber Warriors in the 21st Century*. (Springer, 2024)
104. Badalić, V.: *Preventive Warfare: Hegemony, Power, and the Reconceptualization of War*. (Springer Nature, 2024)
105. Jonsson, O., Käihkö, I.: *Non-Military Warfare: A War of Our Time*. Taylor & Francis (2025)
106. Ziehr, S., Merkt, P.H.: Strategic resilience in human performance in the context of science and education - perspective. *Front. Psych.* **15**, 1410296 (2024). <https://doi.org/10.3389/fpsy.2024.1410296>
107. Marquis, S.L.: *Unconventional warfare: Rebuilding US Special Operations Forces* (1997)
108. Freudenberg, D.: Irreguläre Kräfte und der interessierte Dritte im modernen Kleinkrieg. In: *Die Komplexität der Kriege*, Springer, 179–187 (2010)
109. Souchon, L.: *Strategy in the 21st Century: The Continuing Relevance of Carl von Clausewitz*. (Springer Cham, 2020)
110. Ohlhorst, D., Schön, S.: Constellation analysis as a means of interdisciplinary innovation research-theory formation from the bottom up. *Historical Social Research/Historische Sozialforschung*. 258–278 (2015)
111. Merkt, P.H., Ziehr, S., Voigt, T., Bickelmayer, J., Toursarkissian, M.: *Bedarf und Konzept eines einheitlichen Curriculums für den Master Medic. Taktik+Medizin* (2025)
112. Diakopoulos, N., Friedler, S., Arenas, M., Barocas, S., Hay, M., Howe, B., Jagadish, H.V., Unsworth, K., Sahuguet, A., Venkatasubramanian, S., Wilson, C., Yu, C., Zevenbergen, B.: Principles for accountable algorithms and a social impact statement for algorithms. Online guideline document. <https://www.fatml.org/resources/principles-for-accountable-algorithms>, FAT/ML (Fairness, Accountability, and Transparency in Machine Learning) (2017)

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.