



Kingdom of Morocco

Contribution under the theme:

**“Application of Artificial Intelligence
in the Military Domain:
Opportunities and Challenges
(Excluding LAWS) in a Context
of International Peace and Security”**

Table of contents

I. Introduction	4
II. Opportunities of AI in the Military Domain (Excluding LAWS)	4
A. Functional Domains.....	4
1. Intelligence, Surveillance, and Reconnaissance (ISR)	5
2. Logistics and Operational Planning	5
3. Decision Support.....	6
4. Cybersecurity and Protection of Critical Digital Systems.....	7
5. Training and Simulation	8
B. Use Cases.....	9
1. Conflict Prevention and Diplomacy.....	9
2. Peacekeeping and Stabilization Missions	10
3. Disaster Management.....	10
4. Risk Mapping According to Context	11
5. Migration Flow Management.....	12
6. Border Surveillance.....	13
III. Challenges of AI in the Military Domain (Excluding LAWS)	14
1. Technical Challenges	14
a. Algorithmic Opacity and Cognitive Biases	14
b. Technological Dependence.....	15
c. Vulnerability to Adversarial Cyberattacks.....	15
d. Lack of Human and Infrastructural Capabilities	16
e. Availability and Quality of Data.....	17
2. Ethical Challenges	19
a. Primacy of the Human Factor and Ethics.....	19
b. Discrimination and Societal Biases.....	19
c. Inexplicability of Algorithmic Decisions.....	20
d. Exacerbation of Disinformation Campaigns	21
e. Responsibility and Accountability.....	22
3. Strategic Challenges	22
a. Risks of Overconfidence and Strategic Errors	22
b. Accessibility to Non-State Actors and Terrorist Groups.....	24
4. Legal Challenges.....	25
IV. Conclusion.....	25
References	26

I. Introduction

The rapid emergence of artificial intelligence (AI) is profoundly transforming military strategies and capabilities across the globe. The international community is closely monitoring this technological revolution with immense potential, while also expressing concern over the risks it entails.

This transformation, pushed by AI, presents significant opportunities for conflict prevention, civilian protection, and support for peace operations. However, alongside these hopes, AI raises serious ethical, legal, and security concerns. In the absence of adequate human oversight, AI could escape our control and threaten our societies. Moreover, military uses of AI have primarily drawn attention through the lens of Lethal Autonomous Weapon Systems (LAWS), often referred to as “*killer robots*”, which raise unprecedented questions about delegating life-or-death decisions to machines.

Aware of these issues, the international community is actively debating whether to regulate or ban such systems. Nevertheless, the scope of this document is deliberately limited to AI applications in the military domain excluding LAWS, in order to examine more broadly how AI can support military missions while respecting international peace and security.

This approach aligns with the growing global interest in military AI beyond autonomous weapons alone. On December 24, 2024, the UN General Assembly adopted the first Resolution on military AI, entitled “*Artificial Intelligence in the Military Domain and its Implications for International Peace and Security*”. Other processes have also emerged in this framework, such as the first-ever International Conference on the Role of AI in the Implementation of the Chemical Weapons Convention (*October 2024, Rabat, Morocco*), the Summits “REAIM - Responsible AI in the Military Domain” (*September 2024, Seoul, Republic of Korea ; February 2023, The Hague, Netherlands*), the Political Declaration on the Responsible Military Use of AI and Autonomy (*launched at the REAIM Summit 2023, The Hague, Netherlands*) and the Paris Declaration on Maintaining Human Control in AI-Based Weapon Systems (*AI for Action Summit, February 2025, Paris, France*).

This document explores the opportunities and challenges related to the application of AI in the military domain (excluding LAWS). It draws primarily on UN studies and reports, notably from UNIDIR (United Nations Institute for Disarmament Research) [1], UNODA (Office for Disarmament Affairs), and other relevant agencies, in order to highlight best practices, concrete examples, and related ethical concerns.

Following this introduction, Section II will present the main areas of positive military AI applications as well as use cases ranging from conflict prevention to peacekeeping and disaster response operations. Section III will examine the technical, ethical, strategic, and legal challenges that must be addressed to ensure responsible use of these technologies in line with the law of international relations, particularly international humanitarian law (IHL). Finally, the conclusion will offer a summary of key findings and recommendations, in a spirit aligned with the principles of the UN Charter and multilateral cooperation.

II. Opportunities of AI in the Military Domain (Excluding LAWS)

AI application to the military sector is not limited to autonomous weapons. It encompasses a much broader range of “upstream” functions (intelligence, logistics, information analysis, etc.) that can enhance operational efficiency while reducing risks for both military personnel and civilians. In this second section, we distinguish between the main areas of opportunity AI offers to improve non-lethal military capabilities and concrete use cases illustrating how these technologies can serve the goals of peace and security.

A. Functional Domains

1. Intelligence, Surveillance, and Reconnaissance (ISR)

Intelligence, Surveillance, and Reconnaissance (ISR) form a fundamental pillar of military operations, enabling the collection and analysis of information about the environment, enemy forces, and potential threats. AI offers unprecedented opportunities to enhance ISR capabilities by automating the analysis of vast data streams and detecting patterns that a human eye might miss. For example, computer vision algorithms [2] can review real-time aerial or satellite imagery and identify areas of interest (vehicles, infrastructure, troop movements) with increased speed and accuracy. When coupled with surveillance drones or ground sensors, AI enables continuous coverage of large areas, providing military analysts with a more comprehensive and up-to-date “situational awareness” [3].

The AI embedded in UAVs (*Unmanned Aerial Vehicles*) can help process video feeds and differentiate, for example, an innocuous civilian gathering from a potentially hostile assembly, offering hope of providing early warnings to authorities to defuse likely crises. In other words, AI’s predictive analysis of ISR data can contribute to the early detection of conflicts (*early warning*).

A major advantage of AI in intelligence is its ability to merge and correlate heterogeneous data. It can aggregate information from multiple sources (satellite imagery, communication interceptions, open-source intelligence from social media) and extract actionable intelligence. For example, AI systems can scan social media and detect a sudden rise in violent discourse or false rumors in a particular region, signaling an increase in tensions. Combined with physical terrain surveillance data, these analyses can provide a holistic view of the security situation.

In operational settings, AI also enhances the speed of intelligence processing and dissemination. Where an analyst might take hours to analyze a set of images or reports, an algorithm can perform the task in a few seconds, enabling commanders to obtain a more up-to-date picture of the operational theater. This acceleration of decision-making time through AI can prove decisive in maintaining initiative on the ground. However, it is important to note that this speed presents challenges, as military decisions based on automated analyses must be carefully verified to avoid errors.

In summary, in the ISR domain, AI represents an opportunity to multiply the “eyes and ears” of the military. Whether it is monitoring vast borders, tracking terrorist group movements, or keeping an eye on isolated conflict zones, intelligent sensors can provide constant surveillance where human resources would be limited.

2. Logistics and Operational Planning

First, AI can improve strategic and tactical planning for operations. By simulating different logistical scenarios, algorithms can assist planners in developing the optimal deployment of resources. For instance, AI systems can optimize convoy routes based on multiple constraints (road security, weather conditions, fuel consumption) and quickly recalculate new routes if an unexpected situation arises.

Similarly, in preparing for an operation, AI can propose alternative arrangements for forces and supplies, detecting deviations from the initial plans and suggesting proactive corrections. AI-assisted scenario planning, through tools such as virtual wargaming, allows decision-makers to explore multiple hypotheses to mitigate unforeseen circumstances [4].

Second, AI contributes to making the supply chain more responsive and efficient. Many military organizations are faced with an increasing volume of data to manage: computerized inventories, vehicle tracking sensors, equipment health status reports, etc. An algorithm can continuously process these information flows and trigger alerts or automatic decisions.

In complex international deployment contexts, such as UN missions or multinational coalitions, AI can also facilitate logistical coordination between partners. It can act as a conduit to share information about available resources, urgent needs, and propose mutualizations.

For example, in a peacekeeping mission where several contributing countries provide units, an AI platform could indicate that Contingent A has a surplus of medical supplies that could be redeployed to Contingent B, which is facing a health crisis in its area. By improving global visibility, AI strengthens logistical solidarity among allies and optimizes collective resource allocation [5].

Finally, automation enabled by AI also extends to logistical robotics. Autonomous cargo transport vehicles (on land, or cargo drones in the air) started to be tested to deliver supplies in dangerous areas without human risk.

Although still emerging, these systems use AI to navigate uncertain terrain, avoid obstacles or threats, and reliably reach their destination. The ability to send an autonomous supply robot to resupply an advanced position under enemy fire, or evacuate injured personnel in an autonomous vehicle, would represent a major humanitarian gain by eliminating the exposure of logistics soldiers to ambushes.

In sum, AI brings precision, agility, and foresight to logistics and operational planning. Better forecasting reduces *the fog of war* regarding resource status, while algorithmic optimization enhances the resilience of armed forces in the face of unexpected events.

3. Decision Support

In the command and control (C2) process, AI is seen as an asset to assist decision-makers, whether it is an officer in the field who needs to react in seconds or a staff working on a complex strategy. Decision support via AI takes the form of systems capable of synthesizing information, proposing action options, and assessing the likely consequences of each choice, all in support (not replacement) of human command [6,7].

One of the key applications is real-time decision support systems. In the face of an evolving tactical situation (e.g., the unexpected advance of an enemy column or a sudden deterioration of weather affecting helicopters), an AI tool can instantly aggregate relevant ISR data, compare the situation to similar recorded cases in a database, and provide the commander with several response scenarios with their estimated benefits and risks. The officer still retains the final decision, but benefits from *accelerated analysis* and an *objective evaluation* provided by the machine. This can reduce the risk of a decision made under panic or cognitive bias by offering an informational counterbalance. For example, if AI indicates that a temporary strategic withdrawal has a 90% chance of preserving the unit intact while awaiting reinforcements, whereas the instinct of the moment would push to hold the position, decision-makers can reconsider their initial response in light of this data [9].

Another aspect is high-level predictive analysis. AI models can be trained on vast datasets (historical conflicts, military maneuvers, political and economic indicators) to identify trends that may influence strategic decisions. For instance, AI could be used to assess the likelihood of various regional crisis scenarios occurring within the next five years, integrating variables such as demographics, ethnic tensions, climate change effects, etc. Such a tool does not predict the future with certainty, but it quantifies risks and draws decision-makers' attention to factors that might otherwise be overlooked.

In national military contexts, AI-augmented command centers are emerging, where virtual assistants are part of the planning team. These assistants can respond to natural language queries, such as “*What is the probability of an enemy counterattack if we take city X?*” by providing a structured analysis based on available data. They can also autonomously monitor the progress of an operation and immediately alert leaders to anomalies or opportunities. In

essence, they act as *informational guardian angels*, ready to warn or advise, while relieving human analysts of tedious tasks, such as going through hundreds of incident reports to extract a trend.

It is important to note that in this vision of augmented intelligence, AI is not meant to replace human judgment, but to enhance it. Ethical values and common sense in command remain paramount. AI only provides factual elements and probable inferences.

By leveraging AI for decision support, armed forces also hope to reduce errors and minimize collateral damage. Better informed and advised, a military leader will be less likely to make a hasty decision that could put civilians at risk or compromise the mission. For example, before launching a night operation in a village, an AI system might alert: “*High probability of civilian presence at this time based on mobile phone data, adjust the assault timing*”. This type of contribution helps integrate international humanitarian law into the decision-making process itself.

Ultimately, AI-based decision support is an opportunity to address the growing complexity of modern conflicts. On hybrid battlefields, where military actions, cyberattacks, information warfare, and public opinion pressure are intertwined, decision-makers are overwhelmed with data and constraints.

4. Cybersecurity and Protection of Critical Digital Systems

In the digital age, cybersecurity has become as vital a military concern as the physical defense of territory. Armed forces, defense infrastructures, and even field equipment heavily rely on computer networks and software, which can be targets for cyberattacks. AI plays a dual role here: strengthening the security of critical systems against sophisticated threats and helping to counterattacks carried out by malicious actors in cyberspace.

First, AI enhances intrusion detection and threat detection capabilities (AI-powered IDS) in networks [8,9]. AI systems with machine learning can establish a real-time baseline of normal operation for a military computer network (regular data flows, legitimate processes underway, authorized connections), then detect any suspicious deviation that could indicate an attack. For example, if an unknown malware begins to spread on the logistics network, AI will spot abnormal signatures or behaviors (like a series of unusual connections or unauthorized code execution) and can immediately alert human operators or even initiate threat isolation measures. This increased responsiveness is crucial in the face of attacks that sometimes unfold in seconds. Similarly, AI excels at analyzing large amounts of computer logs to find weak signs of a slow and stealthy intrusion campaign that a human analyst, overwhelmed with data, might miss [10,11].

National critical infrastructures, such as power plants, transportation networks, and strategic communications, can also be protected by AI. For example, a smart grid could use AI to distinguish a fluctuation due to a technical failure from one possibly caused by a cyberattack attempting to trigger a failure. In the event of suspected attack, the system can initiate backup procedures, inform military command, and isolate the compromised part of the network. Moreover, AI systems can simulate the impact of a cyberattack on an entire infrastructure to help authorities plan countermeasures and redundancies. This cyber resilience is particularly crucial in developing countries, where infrastructures are sometimes less robust.

On the offensive side, even though this document focuses on defensive and support uses, it should be noted that AI can be used to counter disinformation campaigns carried out by adversaries, which impact overall security. NLP (Natural Language Processing) algorithms can detect the coordinated spread of fake news, incitement to hate, or online propaganda, often orchestrated by bots. AI provides means to monitor these maneuvers and respond with counter-narratives, helping protect local populations from manipulations that could incite violence and create tension hotspots.

In sum, securing military cyberspace with AI has become indispensable. AI-powered cyberattacks are already targeting critical infrastructures and even peace operations. In response, AI on the defender's side serves as an adaptive shield: the more attacks evolve, the better it learns to detect them.

5. Training and Simulation

Soldiers train in an immersive virtual environment, illustrating how AI and simulation enhance operational readiness while reducing costs and risks. Traditionally, preparing troops for combat or crisis situations relies on intensive training, tactical exercises, and maneuver simulations. AI is revolutionizing this field by enabling smart virtual trainers and realistic simulation scenarios at a level never before achieved. The goal is twofold: to improve the effectiveness of military training while saving resources (ammunition, fuel) and avoiding the inherent dangers of real-life exercises.

One of the notable advances is the development of immersive training simulators using virtual reality (VR) or augmented reality, combined with AI engines to populate and animate the virtual environment. For example, a group of soldiers can operate in a simulated VR battlefield, confronted by enemies controlled by AI that react credibly to the trainees' actions. These virtual enemies do not simply follow predefined behaviors: thanks to Reinforcement Learning, they can adapt their tactics, learn from the player's mistakes, and present a renewed challenge in each session. This means that a soldier or unit in training will not simply repeat a fixed scenario but can be surprised by unforeseen ambushes or complex tricks from the AI, just like in reality. As a result, the learning curve is heightened, as trainees must refine their reflexes and coordination to overcome an evolving virtual opponent.

Human instructors also benefit from AI in tracking and evaluating performance. Systems can automatically analyze the actions taken during the exercise (movements, reaction time, shooting accuracy, adherence to communication protocols) and generate detailed feedback. For example, after a peacekeeping mission simulation in an urban environment, AI could report that the team failed to monitor a blind spot for 30 seconds, or that radio communication was saturated with redundant information—points for improvement that the instructor could address. This objectivity and granularity of evaluations help identify subtle weaknesses that might be hard to detect with the naked eye.

Another benefit is the personalization of training. Thanks to AI, scenarios can be adapted in real time to the trainee's level, much like a video game adjusts its difficulty. If a platoon excels in one area, AI can increase the complexity (for example, by introducing additional random events, such as a sudden equipment failure in the scenario) to push them to perform better. Conversely, if a helicopter pilot trainee struggles with landing maneuvers, the smart simulator can offer more targeted exercises in that area until mastery is achieved. This is sometimes referred to as "smart tutoring," where AI acts as a virtual coach that adjusts the program based on individual progress.

Finally, beyond soldiers, AI is also used to train decision-makers and planners through strategic simulations. AI-assisted war games allow military leaders to test campaign plans in a risk-free environment. AI can play the role of all actors (enemy, allies, civilian actors) and provide rapid feedback on the consequences of a strategy; for example, identifying that a particular deployment plan would lead to a logistical deadlock in X days. These virtual lessons, acquired before real action, are invaluable for avoiding costly mistakes in terms of human lives.

In conclusion, AI is redefining military training as a virtual laboratory where one can learn from mistakes without paying the bloody price, and refine skills in various contexts at will. It is an opportunity to prepare more versatile forces, better prepared for surprises, and therefore more capable of completing their missions with success and humanity.

Naturally, this requires investments in simulation equipment, the development of faithful storyline content, and training instructors in new methods. But these initial costs are more than offset by the savings made (less real ammunition used in training, fewer accidents during real maneuvers) and by the elevation of the operational level of the forces.

B. Use Cases

After reviewing the key functional areas where AI offers benefits for armed forces, let us now look at specific use cases of these technologies in contexts related to peace and international security. The goal is to concretely illustrate how AI can be put to use for objectives such as conflict prevention, support for peace operations, disaster response, and migration management. Each subsection discusses a particular scenario or area where AI brings innovative solutions. These examples will also highlight the cross-cutting nature of the previously described opportunities: we will see those capabilities in ISR, decision support, data analysis, etc., combine in practice to address complex challenges on the ground.

1. Conflict Prevention and Diplomacy

Preventing a conflict before it erupts is arguably the noblest and most difficult mission of international diplomacy. The *preventive diplomacy* seeks to identify early warning signs of crises and promote dialogue to defuse tensions. Artificial intelligence becomes a valuable ally in this effort by providing tools to detect latent conflict dynamics and even facilitate mediation between stakeholders.

A first concrete contribution of AI is the automated detection of weak conflict signals. Big data analytics models process enormous amounts of heterogeneous data—NGO reports, economic data (unemployment, inflation), weather events, social media exchanges, local violence incidents—to spot combinations of factors that have historically preceded crises.

For example, a sudden rise in the price of staple goods combined with increasingly polarized rhetoric on social media in a region where different ethnic communities coexist might be flagged by AI as an alert.

Such AI-assisted early warning systems could be a topic of studies, and some pilot programs have shown that they can anticipate outbreaks of violence by several months, giving diplomats a window of opportunity to intervene.

Of course, AI is not infallible in its predictions, but it allows for much faster processing of information and brings attention to areas that traditional human analysis might not have prioritized.

2. Peacekeeping and Stabilization Missions

UN peacekeeping operations, deployed in fragile post-conflict environments, face significant challenges: civilian protection, ceasefire monitoring, support for stabilization and reconstruction, often in large areas with little infrastructure. The introduction of AI in these missions opens new prospects for enhancing the effectiveness of peacekeepers and improving the security of both local populations and UN personnel on the ground.

Peace missions are increasingly moving toward an integrated civil-military approach, where understanding the local context is crucial. AI can assist the civil component of the mission in conducting rapid socio-economic analyses to guide stabilization projects. For example, by cross-referencing demographic data, the history of violence, and complaints gathered from local populations, an algorithm could indicate which areas are most vulnerable to conflict resurgence due to insufficient public services. Mission leaders can then advocate to governments or donors

for prioritizing these areas for development aid. Thus, AI indirectly helps to strengthen peace by more effectively targeting reconstruction efforts.

Another concrete use case concerns the security of UN personnel. Unfortunately, peacekeepers and humanitarian workers are sometimes targeted. AI can be used to assess the risk of attack against the mission by analyzing online activity from hostile groups, detecting patterns around UN bases (armed reconnaissance, enemy drones), or anticipating unrest (orchestrated protests against the mission via social media). With better information, the mission can adapt its protective posture, reinforce certain positions, or limit non-essential movements during high-risk periods. This can be seen as a form of a “virtual blue helmet” that constantly watches over the mission.

Finally, AI can improve the internal management of the mission. For example, by optimizing patrols: an algorithm can suggest patrol patterns that more efficiently cover a given area based on past incidents and community feedback. Or, manage information: a peace mission generates numerous reports, summaries, and analysis notes every day. An AI assistant can filter and synthesize this information for mission leadership, highlighting critical points and preventing important alerts from getting lost in the mass of documents. This aligns with the goal of digital modernization of peace operations [12].

In summary, in peacekeeping and stabilization missions, AI acts as a force multiplier for “blue” forces: it extends the reach of peacekeepers’ eyes and ears, accelerates the flow of intelligence, anticipates threats, and helps win hearts and minds by adapting efforts to local needs. Of course, all of this must be done responsibly: the UN will have to ensure that these technologies are used in line with its mandate and international law. For example, drone surveillance is carried out without infringing on national sovereignty, as data collected is shared with the concerned authorities to build trust. Similarly, any personal data processed by AI (expressed opinions, etc.) is protected to avoid stigmatizing individuals.

3. Disaster Management

The armed forces are often mobilized in support during natural disasters or major humanitarian crises, providing logistics, rescue operations, and coordination. Artificial intelligence, with its capacity for rapid and predictive analysis, proves invaluable at all stages of disaster management: before (prevention and preparation), during (emergency response), and after (recovery). For the military and civil protection forces, it offers tools to save lives more effectively and target aid more accurately.

In the prevention/preparation phase, AI can help map at-risk areas. For example, by combining climatological, topographical, and land-use data, predictive models anticipate the areas most likely to be flooded during the next rainy season or coastal locations vulnerable to storm surges. This information, provided to military and civilian authorities, allows for strengthening dikes, planning potential evacuations, or pre-positioning emergency supplies (tents, food) near threatened populations.

During the acute phase of a disaster, AI enhances the capacity for reaction and coordination. Take, for example, a major earthquake striking a region. Immediately, drones fly over the area to assess the damage. In real time, computer vision algorithms interpret the images: they identify collapsed buildings, cut-off roads, and crowds of people in distress, with precise geolocation. In a few hours, a map of the destruction and priority needs can be drawn, where it would previously have taken days of human assessment. Military and civilian rescue teams can then target their efforts: sending civil protection units urgently to areas where a collapsed building still holds pockets of survivors, dispatching helicopters to isolated neighborhoods blocked by debris, etc. AI can also analyze emergency communications: for instance, sorting through thousands of social media messages or calls for help to identify those mentioning critical situations (fires, trapped individuals) and reporting them to operators.

In a complex humanitarian crisis (due to conflict, famine), AI can help military forces manage the logistics of aid distribution. It optimizes convoy routes (as seen in the logistics section) but also monitors the efficiency of distribution. For example, by analyzing the mobility data of displaced populations via anonymized mobile signals, AI could estimate if a refugee camp is receiving a sudden influx of people, indicating that another area is under-aided, which allows for balancing efforts. Furthermore, AI can predict the domino effects of a disaster: if a logistical bridge is destroyed, which areas will become isolated and how long before they need an alternative supply? This kind of anticipation is crucial to avoid secondary crises, such as epidemics due to a lack of access to drinking water.

In the recovery phase, after the immediate emergency, AI continues to be useful. For example, in evaluating the damage exhaustively. Post-catastrophe satellite images processed by AI can provide an accurate assessment: X bridges destroyed, Y kilometers of roads unusable, Z homes razed. This helps plan reconstruction and estimate financial needs, crucial elements for mobilizing the international community. Additionally, AI can help rebuild better by simulating, through models, the effectiveness of different architectures to withstand future tremors, thus guiding reconstruction choices (materials, standards).

In conclusion, AI in disaster management allows for valuable time savings, optimal allocation of limited resources, and the avoidance of secondary crises. It provides decision-makers with an almost instantaneous overview during a crisis, where uncertainty once prevailed. Of course, AI does not replace human mobilization and solidarity, but it amplifies them. Rescuers still need to go on the ground, but they do so more effectively, guided by AI analysis.

4. Risk Mapping According to Context

Each geopolitical context presents a unique set of risks and threats. Risk mapping involves identifying, spatializing, and assessing these potential threats to better protect against them. AI proves particularly effective in establishing and updating these dynamic maps, taking into account a multitude of complex factors. Whether in a conflict setting, transnational tensions, or organized crime, AI helps military planners and authorities see on the map what is invisible to the naked eye: emerging risk zones.

In the domain of transnational crime (arms trafficking, drug trade, maritime piracy), AI is used to map probable routes and rear bases. For example, at sea, analyzing the trajectories of ships combined with anomaly detection (like a cargo ship turning off its AIS transponder) allows for the identification of areas where illicit transfers may occur in international waters. A map of suspected transshipment points can be produced, directing national navies or naval missions to increase surveillance in these locations. Similarly, on land, algorithms identify patterns in smuggling routes (for example, a mountain pass that sees abnormally frequent movements at night indicating cross-border trafficking).

A concrete and interesting case is the mapping of risks related to intercommunal conflict in a country under tension. Imagine a fictional country with multi-ethnic provinces where stability is fragile. By analyzing socio-economic indicators (unemployment, land access), expressed grievances (via local media), and recent violent incidents, AI could produce a map showing which provinces or districts are on the brink of explosion.

Managing migration flows is another area where mapping risk is crucial. For example, AI systems could be used to assess which migration routes might see a sudden influx (due to conflict or a policy change). Having this mapped out strategically allows for preparation in terms of reception or control, and for countries of origin or transit, to anticipate the internal impact (informal settlements, potential tensions with local communities).

What makes AI powerful for risk mapping is its ability to take into account complex correlations in space and time. It can reveal, for example, that *“when the dry season peaks and access to water drops in a region, the risk of conflict between herders and farmers rises by 80% within*

the next two months”. Once this spatiotemporal correlation is identified, it allows a window of risk to appear on the map with a timer. Authorities can then take preventive actions (water distribution, inter-community dialogue) precisely during these windows. Without AI, such correlations could remain unnoticed or discovered too late.

Of course, a map is just a snapshot at a given moment; the advantage of AI is its ability to update it continuously, like a living map. This capability is a paradigm shift: moving from static annual risk reports to a near-real-time dashboard. This requires good integration of information systems and inter-agency sharing (intelligence, armed forces, civil protection, etc.), which in itself is an organizational challenge.

In summary, contextual risk mapping through AI offers decision-makers a data-driven, comprehensive view, allowing them to be proactive rather than reactive. It is, in a way, a modern extension of traditional strategic mapping, but enriched by machine intelligence. As long as these maps are interpreted with discernment (AI can make mistakes or overlook intangible factors), they provide a valuable tool for security, diplomatic, and humanitarian planning.

5. Migration Flow Management

Mass migrations, whether caused by conflicts, instability, or disasters, pose a major challenge to regional stability and human security. Managing migration flows involves saving lives, ensuring dignified reception of refugees, while preventing infiltrations of fighters and maintaining social cohesion in arrival areas. AI can help better understand, forecast, and manage these complex population movements, complementing diplomatic and humanitarian efforts.

A key contribution of AI is the prediction of migration flows. By identifying the determinants pushing populations to leave (violence, economic collapse, drought), predictive models can estimate the magnitude and direction of future migrations. For instance, by monitoring indicators in a fragile country (intensifying fighting in a province, catastrophic harvests signaled by satellite imagery), AI could alert that an exodus of a group of people to neighboring borders is likely in the coming weeks. From a security perspective, destination countries can also anticipate pressure on their borders and coordinate with UN agencies to respond, balancing control with assistance.

AI also contributes to real-time detection of movements. Systems combining satellite images, border sensors, and mobile data can spot columns of migrants on the move or sudden influxes at border posts. This allows for rapid rescue operations (e.g., dispatching units to assist a group stranded in the desert before they die from thirst). Of course, this must be done cautiously to avoid giving the impression of intrusive surveillance of vulnerable populations, with the primary aim being humanitarian.

On a more structural level, AI can help optimize the management of camps and migration routes. By analyzing data on the length of stay in transit centers, the rate of deportation or resettlement, AI can recommend improvements: certain temporary camps are becoming overcrowded, so it is better to open another one in anticipation; certain administrative processes create bottlenecks causing crowds to form.

An interesting angle is the fight against human traffickers and smugglers who exploit migration flows. Criminal networks often take advantage of migrants’ desperation, organizing dangerous crossings, and even human trafficking. AI can analyze intercepted communications or financial patterns to help dismantle these networks. For instance, by identifying on social media posts targeting would-be migrants with false promises of passage, and assisting law enforcement in tracing recruiters and smugglers. This contributes to the security of the migrants themselves and stability, as these criminal activities fuel corruption and violence.

In summary, managing migration flows through AI is an exercise in balancing humanity and security. AI offers a humanitarian radar to see ahead and better organize reception, benefiting both migrants (dignified reception conditions) and host or transit countries (less chaos, better resource allocation). It also provides a safety filter to isolate dangerous elements hiding within these flows, without criminalizing the masses. Used wisely, AI can help transform a phenomenon often managed in urgency into a more governed, predictable, and cooperative process, thus reducing international tensions around the migration issue.

6. Border Surveillance

Border control and surveillance are essential attributes of state sovereignty, but they are becoming increasingly challenging in the face of modern cross-border threats: infiltration by armed groups, illicit trafficking, irregular migration, etc. Borders, sometimes stretching over thousands of kilometers across diverse terrains (deserts, mountains, forests), cannot be effectively monitored by traditional human means alone [13-14]. This is where AI comes in, at the heart of what are often called “smart borders” [13].

AI, coupled with sensor networks, enables constant vigilance along borders. Motion-detecting cameras with vision algorithms can distinguish between an animal, a vehicle, or a human crossing the border line, both day and night, and immediately alert the guard posts in case of suspicious intrusion.

Beyond detection, AI is useful for anticipating attempted crossings. By analyzing trends (time of year, favorite crossing points) and combining with intelligence information, software can pinpoint the most likely border segments for an upcoming attempt. For example, noting that each year before New Year’s, there is an increase in clandestine crossings in a particular wooded area. With this knowledge, the command can strengthen AI/human monitoring during this critical period.

On maritime borders, AI is just as crucial. Coastguards use radar and imaging systems to monitor the coasts; AI can filter out the noise of the sea and detect small boats often undetectable by traditional radar. The Frontex agency in Europe is testing algorithms to spot boat people in the Mediterranean via drones or satellites, to guide rescue ships [15]. Similarly, for combating illegal fishing or the infiltration of commandos by sea, AI helps discriminate abnormal vessels among thousands of marine signals.

Another aspect is the automated management of official border posts. AI systems for facial recognition and behavioral analysis can speed up controls while spotting potentially suspicious individuals (e.g., someone with a wanted notice or a security alert).

These technologies, already present in some airports, raise privacy concerns but can increase efficiency. At a very busy land border post (e.g., between two neighboring countries with heavy traffic), AI can streamline regular traffic (fast lanes for known trucks, etc.) while allocating human attention to atypical or risky cases.

Where AI also excels is in merging border data. For example, aggregating information from an underground seismic sensor (detecting footsteps or vehicles), an aerial drone, and a ground patrol to create a unified picture of an intrusion attempt. Rather than each source operating separately, AI combines and reduces the informational noise, issuing an alert only when the entire system converges on a significant event. This avoids repeated false alarms that could otherwise saturate guards and lead them to error.

A sophisticated use case is that of the virtual border. For example, in mountainous areas without physical fencing, one could imagine a “boundary” defined by an invisible network of sensors and drones. AI manages this network: one drone automatically follows a target spotted by a ground sensor, while another goes to recharge, etc. This creates a kind of active, modular barrier.

The goal is not to militarize borders excessively, but to multiply detection capabilities without requiring a human presence everywhere.

In the end, AI redefines the border not as a simple static defensive line, but as an intelligent space where information circulates, and detection-prevention takes precedence over late reaction. This improves security while potentially reducing the use of force (as one can intercept early, without confrontation, or deter via visible vigilance).

After this detailed exploration of AI's opportunities and use cases in the military domain outside of lethal autonomous weapon systems (LAWS), it is clear that these technologies can play a transformative and positive role. From intelligence to logistics, from conflict prevention to humanitarian crisis management, AI offers decision-makers and responders new means to act more efficiently and with greater anticipation. However, realizing these promises is not automatic: it depends on wise choices, ethical oversight, and careful risk management. That is why the next section will be dedicated to the challenges—technical, ethical, strategic, and legal—that must be overcome to ensure responsible integration of AI in the military domain.

III. Challenges of AI in the Military Domain (Excluding LAWS)

Despite the clear advantages of artificial intelligence (AI) in the military domain, its integration raises numerous challenges that must be identified and addressed. These challenges vary in nature and are often interconnected: technical (system reliability, cybersecurity vulnerabilities, etc.), ethical (human respect, bias, privacy), strategic (power imbalances, escalation risks), and legal (normative gaps, accountability).

This section examines each of these four major types of challenges, detailing their manifestations and drawing from analyses and recommendations in UN reports. The goal is twofold: to prevent potential abuses related to military AI and to propose strategies to ensure that AI remains a tool for peace rather than an additional risk factor [16].

1. Technical Challenges

a. Algorithmic Opacity and Cognitive Biases

One of the primary technical obstacles to the seamless adoption of AI in military environments is the so-called “black box” issue. Many AI systems, especially those based on deep learning techniques, function in ways that are too complex for humans to easily understand. They can provide recommendations or make decisions (e.g., identifying a target in an image) without making clear *how* they arrived at their conclusions. This algorithmic opacity poses a significant issue in a defense context where trust and explainability are essential. Analysts and military leaders naturally struggle to trust analyses they cannot explain or verify through clear reasoning. For instance, imagine a software stating: “*this vehicle is hostile with a 92% probability*”. If it is unclear that this conclusion was due to the system mistakenly identifying a shadow as a weapon (a common error), an action could be taken based on incorrect information [17].

This opacity is often compounded by biases in algorithms. Biases can arise from the training data or the design of the model itself. For example, a facial recognition system trained primarily on specific ethnic background faces may perform poorly on individuals from other ethnic backgrounds, a bias that has been documented in many studies. Translated into military contexts, this could mean that a person identification system performs poorly on certain populations, leading to potentially grave errors (e.g., misidentifying a civilian as a suspect).

Human cognitive biases also interact with AI opacity. Placing excessive trust in automation is a known phenomenon. If an AI system has consistently provided good results, operators may blindly trust its decisions, even when it makes mistakes. This tendency is enhanced if the

algorithm does not provide explanations, making it harder for humans to detect potential errors. For example, during an AI-assisted target identification, military personnel may follow the system's recommendation to fire without questioning it, especially under stress, even though a more detailed human analysis might reveal a mistake (e.g., misclassifying a civilian target). Experts emphasize that human-machine interaction in military contexts is delicate: too much distrust of AI nullifies its benefits, while too much blind trust can lead to disastrous outcomes.

From a technical standpoint, research in explainable AI (XAI) aims to develop models that are more transparent or capable of providing clear justifications for their results. For example, a vision system could state: *"I identified this person as hostile because he/she is wearing a uniform similar to a particular model"*, which would allow for verification of the reasoning. This field should be encouraged in military applications.

Regarding biases, the challenge lies in ensuring diversity in training data and systematically testing models across varied scenarios. Armed forces developing AIs should integrate social scientists or human rights experts into their teams to detect and correct potential biases.

Ultimately, the challenge of opacity and bias is one of reliable intelligibility. To gain the trust of military personnel while preserving ethical principles and non-discrimination, AI must be as *understandable* and *manageable* as possible.

b. Technological Dependence

The widespread introduction of AI into military systems brings with it the risk of increased technological dependence.

From an operational perspective, relying on AI for critical tasks could lead to a degradation of human skills. For example, if logistics planning is entirely optimized by AI software, what happens when the software is unavailable or fails? Will military personnel still be able to develop complex logistics plans "the old-fashioned way"? In military history, we have seen examples of specialized troops losing basic skills (for example, celestial navigation was nearly forgotten in the GPS era, which became problematic if GPS was jammed). A concern is that, by becoming accustomed to AI assistance, forces lose cognitive flexibility and the ability to improvise without these digital crutches. The challenge, therefore, is to maintain a balance between human and machine: AI should be an aid, not a complete substitute. Training should incorporate "manual fallback" scenarios in case technology becomes unavailable [19-20].

Another aspect is dependence on technological infrastructure. Military AIs rely on telecommunications networks, databases, and high-performance computing (HPC) infrastructures. In intense conflict situations, these elements may be destroyed, jammed, or sabotaged. If the entire command and intelligence chain is built around an AI backbone, the neutralization of this backbone by the enemy could paralyze operations. Armed forces must, therefore, plan for degraded modes of operation. This reflects the issue of resilience: ensuring that even without AI, the mission can continue (albeit less efficiently). Moreover, this raises the question of contingency plans in case of a bug: a critical bug in an embedded AI system (e.g., an autonomous vehicle blocked due to a faulty software update) should not endanger lives due to a lack of workaround solutions.

It is recommended to ensure that there is always a "Plan B" without AI or with backup AI systems. By analogy, in civil aviation, although autopilot is ubiquitous, pilots are trained to fly manually if needed, and planes have backup systems. Armed forces should adopt a similar approach (e.g., retaining traditional navigation skills or having backup analog equipment).

c. Vulnerability to Adversarial Cyberattacks

By deploying AI systems, armed forces expose themselves to new forms of attacks from adversaries seeking to deceive or corrupt these systems. The AI itself can become the target of

specific hostile tactics, such as manipulating its sensors, poisoning its data, or hacking its algorithms.

These are called adversarial attacks [21-22], where a malicious actor introduces disruptions designed to deceive the algorithm. For example, painting a vehicle with a particular pattern to make it go unnoticed by a vision system or injecting false data into intelligence channels to mislead the AI.

Researchers have shown that it is possible to deceive a visual recognition AI by presenting it with images altered in ways imperceptible to the human eye, but which the model misclassifies. In a military context, an enemy could exploit this technique to make an autonomous system mistake a legitimate target for an innocuous object, or vice versa.

Connected AI systems (drones, robots, data fusion centers) also present cyberattack surfaces. A sophisticated adversary could attempt to hack the communication network linking the sensors and the AI, either to neutralize it (denial of service) or take control. If a military AI is infiltrated, the enemy could use it to spread false information to commanders, creating chaos and internal distrust.

The robustness of AI systems against these attacks is therefore a major technical challenge. Developers must anticipate adversarial tricks and train models with noisy data or known adversarial examples to immunize them as much as possible.

Moreover, in the field, AIs must be surrounded by traditional cybersecurity countermeasures (encryption of communications, firewalls, intrusion detection) to prevent unauthorized access. A UNIDIR report on AI security [18] emphasizes that the issues of safety and security for automated systems are closely related, and that an AI's adaptability to unforeseen situations must be improved to avoid breakdowns in slightly altered environments.

Meanwhile, the international community is starting to exchange best practices to counter adversarial attacks. Joint exercises among allies, for example, simulate cyberattacks on intelligent military networks to derive lessons. But until global AI governance standards are in place, this technical challenge will persist: every new AI capability must be tested not only under nominal conditions but also under the scrutiny of a clever adversary. This reality requires designing resilient AIs that can continue to function in a degraded mode in hostile environments.

d. Lack of Human and Infrastructural Capabilities

The large-scale adoption of AI in armed forces faces a lack of qualified human resources and appropriate infrastructure in many countries. Designing, deploying, and maintaining effective AI systems requires specialists in computing, data science, and cybersecurity, which not all armed forces have in sufficient numbers.

Additionally, there is a deficit of technological infrastructure in some states: secure data centers, high-speed connectivity across the country, sensors, and digital data collection platforms. Without these prerequisites, the best algorithms in the world cannot be leveraged. For example, if the army of a developing country does not have satellites or drones to collect images, it cannot feed a geospatial analysis AI system; or if networks in an operational area are unstable, an AI requiring continuous updates may not function properly.

Less advanced countries might find themselves unable to keep up. The investment required to implement military AI in developing countries, already facing basic infrastructure challenges, is significant, hence the importance of international cooperation programs: knowledge sharing, regulated technology transfer, helping train local data scientists, etc.

Moreover, even in technologically advanced militaries, it is important to avoid the loss of traditional skills. Soldiers must retain fundamental non-technological skills to cope with potential failures (navigation without GPS, terrain discernment without computer aid, etc.). This requires a dual training effort, which can be burdensome: on the one hand, learning to use new AI tools, and on the other, continuing to exercise “manual” critical thinking.

At the infrastructure level, implementing AI also requires managing data adequately. Many armed forces still lack dedicated data centers or data governance policies (storage, classification, sharing). Without solid infrastructure to collect, clean, and store data, AI cannot be deployed effectively. However, building this infrastructure is expensive and can take time (installing servers, adopting secure clouds, etc.). Public-private partnerships could be a solution, provided the ownership and confidentiality of defense data are well-negotiated.

In summary, the challenge of human and infrastructural capabilities is a call for investment and training. It highlights that the successful introduction of AI is not just a matter of algorithms but requires a profound organizational transformation. Developing national AI plans, including a defense dimension, is encouraged, in order to plan for the long-term development of the necessary skills and infrastructure. This is accompanied by international coordination so that these advances benefit everyone. In this sense, the UN could invest in promoting the exchange of good practices between States.

e. Availability and Quality of Data

The performance of an AI system primarily depends on the quality, quantity, and relevance of the data on which it is trained and that it uses in operations. In the defense sector, however, obtaining reliable and sufficiently abundant data is a challenge in itself.

Firstly, some situations are infrequent or unprecedented, which limits the possibility of training AI. For example, an AI tasked with predicting an adversary’s reactions during a crisis often has only a few historical cases to learn from (each war or diplomatic confrontation is unique). This can lead the algorithm to over-generalize from too few examples.

Similarly, to train an AI to recognize improvised explosive devices (IEDs), thousands of IED images are ideally needed, but each theater produces a limited number of different models. Military personnel sometimes have to resort to data synthesis (artificially generating similar data) to augment training sets, but this can introduce biases.

Secondly, available data may be partial or biased. For instance, if a predictive model of civil unrest is built based on intelligence reports from a particular country, those reports may reflect the biases or interests of that intelligence service, rather than the complete reality. The AI will then learn these biases. The representativeness of the training data is crucial; it must cover the diversity of possible environments and behaviors, otherwise, the AI will fail when it is out of its comfort zone.

Another issue is the shortage of labeled data. Many supervised algorithms require data annotated by humans (for example, hours of video where each frame is labeled as “civilian,” “combatant,” etc.). However, this annotation work is massive and often classified in a military context. Analysts must be mobilized to do it, which takes time and resources. Self-learning (unsupervised) or weakly supervised learning methods are being developed to mitigate this need, but they remain less precise. The availability of reliable data thus becomes a bottleneck: it is not enough to have sensors and archives; these data must be exploitable.

Many developing countries lack usable digitized databases in the security sector. For example, there is no complete registry of past incidents, no detailed digital mapping of terrain, etc. They are essentially starting from scratch to feed AI, unlike countries that have been accumulating digitized intelligence for decades. This backlog is not easily overcome. Hence the importance

of international data-sharing programs when they can be pooled without jeopardizing sovereignty.

Finally, even when data exists, a technical challenge is ensuring its integrity and timeliness. Outdated data can mislead AI (for example, a city plan that does not account for new buildings). This underscores the need for human verification chains and data hygiene: traceability of the source, cross-checking between multiple sources to spot potential inconsistencies (e.g., a sensor reports an event that no other sensor confirms, suggesting a potential decoy).

Furthermore, the lack of high-quality data is a direct impediment to the effectiveness and accuracy of algorithms. Consequently, it is recommended to invest now in defense data management: secure digitization of archives, standardization of report formats to facilitate automated analysis, and the creation of interoperable data lakes between different military branches.

In conclusion, the age-old problem of intelligence—obtaining reliable, up-to-date, and exhaustive information—remains present in the age of AI, simply transposed to the digital field. AI can certainly help fill gaps (by extrapolating trends despite incomplete data), but this has its limits. This technical challenge calls for patience and rigor: taking the time to build solid databases will determine the success of AI in operational situations. Armed forces and the intelligence community must collaborate closely with data specialists to create this foundation, and the United Nations can play a facilitating role by developing standards for the secure sharing of data for peace and security purposes.

2. Ethical Challenges

a. Primacy of the Human Factor and Ethics

At the heart of ethical concerns is the fundamental question of the human role in the use of force and the conduct of war. The principles of international humanitarian law (IHL) and the United Nations Charter are based on the idea that the decision to employ lethal force, particularly, must be made with discernment by responsible human beings. The emergence of AI raises the fear of an erosion of this human primacy. Even outside the extreme case of LAWS, AI can influence decisions regarding engagement, targeting, and neutralization of threats. It is imperative to remind that ethics must guide AI, not the other way around.

AI should remain a tool of assistance and not a total replacement. For example, an AI may prioritize potential targets on a list, but the final decision to strike should be validated by a trained human operator, capable of applying the rules of engagement and making judgment calls (e.g., canceling if in doubt, verifying the proportionality of the attack, etc.). Similarly, in the command chain, one may accept that AI suggests a tactical plan, but it is up to the commander to evaluate its compliance with the law and political objectives before approving it.

Emphasizing human primacy also ensures that fundamental ethical values—respect for life, human dignity, necessity, and proportionality in conflict—are integrated as constraints in the design of systems. For example, an autonomous vehicle patrolling must be programmed to avoid hitting civilians at all costs, even if it means letting an assailant escape. A target sorting algorithm must be capable of *recognizing the abstention*: detecting that there is no clear legitimate target and recommending not to shoot. This involves encoding *ethical rules* into AI (sometimes referred to as “ethical governors”).

Furthermore, the ethical training of personnel becomes all the more crucial in the age of AI. If a military person misunderstands how AI works, they may be tempted to delegate moral responsibility: “the machine decided”. This is unacceptable from the perspective of the law of

war; responsibility cannot be diluted or shifted to a non-human entity. Military personnel must be educated to retain their critical thinking and moral consciousness in the presence of AI. This is to avoid the phenomenon of “*responsibility dilution*” that could otherwise arise (with everyone following automated recommendations without feeling accountable for the final result).

Finally, the primacy of ethics requires setting red lines for the use of AI. For example, the international community could agree that certain decisions should never be made solely by AI, regardless of technical advances. This is partly the subject of the debate on LAWS, where the boundary of autonomy must be defined. Outside of LAWS, we could think of prohibitions such as: no AI to decide on degrading treatments of prisoners (a human must always validate any act on a prisoner, according to the Geneva Conventions), or no AI to assess the “value” of one human life over another (a cold calculation of acceptable collateral damage should not replace human judgment). These ethical discussions must bring together military personnel, lawyers, philosophers, and civil society.

UNESCO produced a Recommendation on AI ethics in 2021, which, while general, sets out principles (transparency, fairness, human oversight, accountability) applicable to the military sector [23]. States could draw inspiration from it for their defense doctrine.

In summary, war is not a mere optimization problem that AI can solve, but a deeply human phenomenon that should remain under human control. The tragedies of past mistakes (such as bombing civilians by mistake, etc.) call not for eliminating humans from the process, but for refocusing them with better decision-making tools. Preserving humanity ultimately means preserving the possibility of compassion, de-escalation, and forgiveness—qualities that no machine will ever possess.

b. Discrimination and Societal Biases

Despite the apparent objectivity of machines, AI can inadvertently reproduce or amplify existing societal discriminations. This is a major ethical issue: ensuring that the adoption of AI does not reinforce inequalities or injustices, whether in peacetime or in the conduct of military operations.

In the context of armed conflicts, algorithmic discrimination can manifest tragically. Imagine an AI that helps choose targets to strike: if, by design, it gives less consideration to the presence of certain categories of people (for example, it poorly detects women and children because training data focused on silhouettes of armed men), it could lead to less careful strikes in areas predominantly inhabited by women and children, resulting in higher collateral damage to these populations—this is a form of discrimination by the tool. Conversely, an AI might overestimate the danger posed by young men of a certain origin at a checkpoint, leading to harsher treatment of them, which violates the principle of equal protection.

Non-discrimination is a fundamental principle of international human rights law. Its respect must be ensured even in algorithms. To achieve this, regular algorithmic audits should be implemented: testing the system with diverse cases, checking the rates of false positives/negatives for different demographic groups, etc. If a bias is detected, the model must be corrected or its use adjusted.

Caution must also be exercised regarding the use of certain characteristics in models. For example, including ethnicity or religion as an input variable in a predictive behavior algorithm would be extremely delicate, as it would officialize discrimination.

However, even without explicitly including them, AI can deduce them indirectly through other data (place of residence, style of dress, etc.). Hence, it is essential to design AI systems based on relevant and legitimate criteria (such as verified criminal records) and to exclude sensitive attributes as much as possible.

Ultimately, ensuring algorithmic justice is a moral imperative that intersects with practical concerns: discriminatory AI undermines public trust and can even provoke unrest. Armed forces must instead appear exemplary in the fair use of technology. This requires transparency, independent audits, and proactive bias correction.

c. Inexplicability of Algorithmic Decisions

When decisions made by or with the help of AI cannot be explained in an understandable manner, it raises an ethical issue of transparency and accountability.

From an ethical standpoint, every act of war or coercion should be justifiable afterward. If a drone strike kills civilians, the reasons for this mistake must be explained to hold accountability and prevent its repetition. However, if the situation analysis is largely automated, one might hear “*it is the algorithm’s fault, we do not know why it did that*”. This situation creates an unacceptable accountability vacuum. Victims or their families are entitled to demand answers (the right to justice and effective remedy). Failing to provide a comprehensible explanation would be a double injustice: not only is there harm, but the clarity about its cause is also denied.

Inexplicability also complicates legal proceedings. Imagine a court investigating a military incident. If an officer says, “the AI system told me it was a legitimate target”, the judge will ask, “on what basis did the system conclude this?” If no one knows (no intelligible software, complex model), how can it be determined whether there was negligence or a violation of international humanitarian law (IHL)? This could hinder the application of the law. In a national context, it would also clash with the principle of democratic accountability: parliaments that oversee executive actions in defense matters must be able to examine and understand the information and decision-making processes that led to a particular operation. If they face a technical wall, democratic oversight weakens.

There is also a moral trust dimension. Military personnel in the field, as well as the public, will struggle to trust a system if they feel it acts in an incomprehensible manner. Trust is essential for ethical acceptance: people are more likely to accept a decision (even a negative one) if they understand the reasons behind it. This is a basic principle of justice: a judge’s decision is reasoned. Similarly, it could be argued that every important decision where AI played a role should be justified in an intelligible manner.

To address this challenge, ethics meets technology: developing explainable AI (XAI) [24]. For example, researchers are working on neural networks capable of providing not only a result but also an “attention map” highlighting which parts of the data most contributed to the decision. In image analysis, this could be visualized by highlighting: “I identified this person as a combatant because I detected the shape of a weapon in their hand”. This type of visual or textual explanation allows an operator to validate or reject the AI’s evaluation. If the algorithm highlights an innocuous object (for example, a shovel confused with a rifle), the human can correct it.

Similarly, hybrid models combining logical rules and machine learning are promoted, where the rule-based part (expressed in “if...then” statements) frames the algorithm and can be easily audited. For example, there could be a non-negotiable rule: “never classify as a target an individual raising their hands or wearing a medical symbol”, coupled with an AI identifying shapes. The explicability here comes from the rule.

International organizations also encourage the creation of algorithmic registers and event logs: each major recommendation or decision made by an AI system should be recorded with its context, so it can be analyzed later by independent experts. Analogous to the “black boxes” of airplanes, there could be “black boxes” for military AIs. In the event of an incident, these logs would be examined to reconstruct the decision-making path. If it is discovered that the

algorithm did X after input Y, researchers could seek to understand why in the lab, even if it was unclear at the time.

Ultimately, lifting the veil on the black box is not only a technical challenge but also an ethical imperative to maintain responsibility and trust. By aligning innovation with transparency, we can ensure that AI remains a tool in service of the law, not a factor of opacity or arbitrariness.

d. Exacerbation of Disinformation Campaigns

AI offers highly effective means to create and spread disinformation, posing a major ethical and security challenge. In modern conflicts, and even below the threshold of war, the information battle is raging. Manipulating public opinion, viral rumors, and false content aimed at causing panic or hatred can dramatically amplify the impacts of a conflict or hinder UN peacekeeping missions. The rise of deepfakes (ultra-realistic videos manipulated by AI) has marked a turning point: we can no longer trust the evidence of our senses to judge the authenticity of information.

Indeed, one can imagine the nightmare scenario of a fake video showing a head of state declaring war or a peacekeeper commander committing atrocities. The immediate reaction could be violent before the fake is exposed. Such artifices exploit emotions and the speed of diffusion on social media to create “information explosions” that are hard to contain.

Ethically, the deliberate use of disinformation is condemnable, particularly when it aims to harm civilians (hate propaganda, etc.). A fully AI-driven cognitive war, aimed at breaking the enemy’s will by flooding its population with fake news, would be a very dark prospect that should be avoided.

On the other hand, AI can be used to reveal disinformation. Algorithms are already scanning networks to detect signs of coordinated campaigns (automated accounts, massive diffusion of the same message, false images identified by pixel analysis or comparison with original image databases). The ethical balance is delicate: disinformation must be countered without sliding into unjustified censorship.

Armed forces and peace missions can use AI to circulate verified information quickly after an incident, cutting off malicious rumors. AI can even help predict the spread of a hoax: which groups are likely to pick it up, in which regions it will spread, allowing for targeted corrective communication.

However, an ethical risk is that public trust in any information, true or false, might diminish in the face of this algorithmic battle. This is referred to as a “post-truth” society where nothing is believed. This is in itself a danger to democracy and peace, as widespread mistrust benefits troublemakers. It is therefore crucial that legitimate actors (States, the UN) maintain a reputation for reliability. If they themselves were caught manipulating information through AI, their credibility would be ruined. Hence the need for an ethical charter: for example, some military forces have publicly committed not to create deepfakes in their communications.

From a legal perspective, emerging discussions are considering assimilating certain severe information attacks to violations of international law (even acts of aggression when they cause damage comparable to armed attacks) [25].

In conclusion, ethics and security converge on this point: preventing a military “infodemic”. AI, a double-edged sword, must be channeled to defend truth rather than spread lies. It is an asymmetric struggle because creating lies attracts more attention than debunking them, but it is a vital cause for international stability. States could cooperate in this sense by sharing hoax-detection algorithms, alerting each other in case of cross-border toxic campaigns, and educating the public (the “cognitive resilience” of populations is also a shield) [26]. Technology must go hand in hand with the elevation of citizen critical thinking, so that AI does not find fertile ground in societies vulnerable to manipulation.

e. Responsibility and Accountability

The introduction of AI in military decision-making complicates the issue of responsibility in case of error or violation. However, in practice, determining who within the state apparatus or the chain of command should be held accountable can become tricky. For example, if a biased algorithm leads to a mistake, should the responsibility fall on the soldier operator (who may have followed their training and could not detect the bias)? The commander who authorized the use of this system? The private AI manufacturer who supplied a faulty tool? The officer who did not adequately supervise the automatic decision? A situation of ambiguity can create what is called an “accountability gap”, a void of responsibility where everyone can pass the blame.

To address this challenge, several approaches are being considered:

- **Decision traceability:** Keeping logs that identify who validated what, at what time. This helps trace responsibility up the chain. For example, if an AI suggested a target and an analyst explicitly confirmed it, the responsibility of the analyst is engaged in addition to that of the commander.
- **Doctrinal clarity:** Establishing in AI operational doctrine that humans are always responsible. Some countries have already integrated into their defense AI principles that AI does not relieve operators and commanders of their responsibility.
- **Review and oversight mechanisms:** Establishing ethics or accountability committees within the military to examine AI-related incidents. Similar to investigative commissions, but specialized. This helps identify whether there was human fault (e.g., failure to adhere to a double-checking procedure) or if the problem was technical (then calling the supplier to fix it, or even holding them contractually responsible).
- **Updating rules of engagement:** These can explicitly outline the boundaries of AI and the points at which human approval is required. Thus, if these rules are violated (e.g., an operator lets AI target without superior authorization), disciplinary responsibility is clearly assigned.

In short, the ethics of responsibility demands that accountability not be diluted in technological complexities. On the contrary, the more technology is present, the more accountability safeguards must be reinforced. For example, one could define that every unit using AI must have an officer specifically tasked with overseeing its use (an “AI control officer”) who will be held accountable in case of issues, much like a security officer is for weapon handling. Such institutional innovations can help keep a human (and responsible) face behind every decision, preventing the machine from becoming a convenient scapegoat or, worse, an uncontrollable actor.

3. Strategic Challenges

a. Risks of Overconfidence and Strategic Errors

The introduction of AI into military planning and operations could lead to excessive confidence in automated predictions or recommendations, potentially resulting in strategic errors with grave consequences. Military history is full of instances where an overreliance on new technology or doctrine led to significant setbacks. With AI, this dynamic might repeat itself if decision-makers overestimate the reliability of their systems.

A first risk is the misinterpretation of AI’s true capabilities. For example, if an army believes that its predictive system can reliably anticipate the enemy’s movements, it might make bold decisions (like splitting forces or launching a preemptive offensive) based on those predictions. If these predictions turn out to be incorrect, the resulting posture becomes dangerous. In 2020, a report highlighted that “*overestimating what AI can accomplish*” can lead to its use beyond

its capabilities and result in major failures [27]. Thus, AI could encourage strategic recklessness if seen as an infallible oracle.

The second risk is related to the automation of decision-making time. AI allows for the acceleration of the “OODA” loop (Observation, Orientation, Decision, Action) to the point where strategic decisions could be made very quickly. While this can be a tactical advantage, at the strategic level, hastiness is often a poor advisor. A system that indicates “retaliate now” in a matter of seconds could create political pressure to act without allowing time for diplomacy or reflection. This is the issue of *compressed reaction time*: fearing being outpaced by the speed of the opposing AI, parties could begin to react almost automatically, risking rapid escalation. At the highest level, this could mean missing opportunities to defuse a crisis through dialogue because the AI “has already acted”.

The increased complexity of systems can also lead to strategic errors. The more automation there is, the more the logistical and operational chain depends on technical conditions. A clever adversary might seek to sabotage AI strategically, for example, by deceiving satellites or hacking a key database, thereby triggering a large-scale erroneous decision (such as committing forces to the wrong location). Such tactics could cause dramatic *miscalculations* if high command does not detect the deception in time.

What can be done to mitigate these risks? First, cultivating a culture of questioning AI. In peacetime, conducting “*red team*” exercises that simulate what happens if AI makes a major error. In times of crisis, do not deactivate critical thinking just because “the machine said so”.

At the international level, a kind of trust measure could emerge: a tacit agreement to keep humans in the loop for strategic decisions, similar to diplomatic “hotlines” to clarify intentions and prevent accidents. For example, partial transparency on the fact that a particular alert system is human-supervised could reassure the adversary.

AI could also be used to promote *stability*: for example, algorithms could be programmed to automatically detect signs of error and block automatic reactions. “Escalation guardians” could be programmed to encourage restraint if a quick decision clearly contradicts diplomatic indicators (e.g., if diplomatic channels are active, it suggests that an agreement is in progress, so do not launch an attack, even if tactical modeling suggests otherwise).

Finally, a strategic awakening might come from recognizing the fallibility of AI. Just as it was accepted that lightning could disrupt a radar, it must always be kept in mind that AI can make mistakes. This precautionary principle should be integrated into war games conducted by General Staffs: testing “what do we do if our AI misleads us?” much like exercises for GPS loss. Having a *manual Plan B* will mitigate overconfidence.

In conclusion, AI does not eliminate the “fog of war”; it merely changes its nature. The risk of surprise or error does not disappear; it may even increase due to excessive digital certainties.

b. Accessibility to Non-State Actors and Terrorist Groups

In contrast to traditional sophisticated weapon systems (ballistic missiles, fighter jets, etc.), AI is widely available as inexpensive commercial software, making it accessible to non-state actors, including criminals and terrorists.

This is a significant strategic challenge: armed groups could use drones equipped with rudimentary AI for targeted attacks, or extremist networks could deploy bots and deepfakes to radicalize and coordinate followers remotely.

There have already been cases of armed quadcopters used by terrorist organizations, and adding autonomous capabilities and computer vision could increase the lethality of these improvised

devices. Similarly, a group could attempt to hack into a national military system using AI tools available on the dark web, threatening collective security.

This democratization of AI disrupts the state's monopoly on legitimate violence. It complicates peacekeeping operations (Blue Helmets may face adversaries equipped with AI for attack or camouflage) and counterterrorism (intelligence services must account for this factor). Strategy and ethics converge here: it is essential to deprive malicious actors of the means to abuse AI.

If the international community agrees on responsible AI usage norms, the deviant behavior of terrorists will appear all the more clearly as criminal and illegitimate, facilitating consensus to combat them.

Moreover, by developing AI countermeasures in advance (defensive AI for filtering propaganda, jamming adversarial AIs, etc.), states can maintain the technological upper hand over these non-state actors.

4. Legal Challenges

From the perspective of international law, military AI is currently evolving within a relative normative vacuum. No universal convention specifically addresses its development or use (ongoing discussions on LAWS under the Convention on Certain Conventional Weapons, which remain limited to lethal autonomy, and the UN General Assembly resolution on artificial intelligence in the military domain and its consequences for international peace and security adopted on December 24, 2024).

This lack of a clear legal framework raises concerns about the possible emergence of divergent, and potentially conflicting, practices without effective means to prohibit them.

Existing law, especially international humanitarian law (IHL), applies to all weapons and methods of warfare, whether or not they involve AI. The principles of distinction, proportionality, and precaution in attack, as well as the prohibition of unnecessary suffering or targeting civilians, remain in force.

However, these rules were conceived in the context of human decision-making, and their interpretation in highly automated scenarios is far from settled. For example, how should the proportionality of a strike be assessed if the damage estimation was based on AI calculations? How far does the duty of precaution extend when it comes to verifying the outputs of an AI before taking action? These are debated questions that the international community continues to address.

IV. Conclusion

Artificial intelligence is poised to profoundly transform the military domain, bringing significant opportunities to better anticipate crises, protect populations, and assist soldiers on the ground. Promising applications are emerging in intelligence, logistics, decision-making, cybersecurity, and training — including in areas such as peacekeeping, disaster management, and the prevention of ethnic conflicts.

However, AI is not a magical solution free of risks. Misused or deployed without safeguards, it could exacerbate tensions, erode humanitarian principles, and generate new threats to global stability.

The technical challenges (system reliability, protection against hacking, data quality), ethical challenges (transparency, absence of bias, accountability), strategic challenges (avoiding uncontrolled escalation, reducing the technological divide), and legal challenges (filling the normative gap and inventing control mechanisms) are immense.

Each of these requires a concerted response from the international community. Faced with cross-cutting and global technologies, it is through cooperation, dialogue, and the establishment of common rules that we can minimize the risks while maximizing the benefits of AI.

In a world already shaken by conflicts, pandemics, and climate change, AI must not become a new source of division or destruction. On the contrary, if collectively mastered, it can help defuse conflicts (through better prevention), reduce mistakes and collateral damage (through enhanced precision and more rational decision support), and strengthen peace operations (through improved logistical efficiency and personnel safety).

In conclusion, the application of artificial intelligence in the military field marks a historical turning point. It brings unprecedented opportunities to build a safer world, where decisions are better informed and where both soldiers and civilians are better protected.

References

- [1] M.A.-É. internationales, undefined 1989, Institut des Nations Unies pour la Recherche sur le Désarmement (UNIDIR). Désarmement: Problèmes relatif à l'espace extra-atmosphérique. New York, UNIDIR, Erudit.Org M Albert Études Internationales, 1989 erudit.Org (n.d.). <https://doi.org/10.7202/702531ar>.
- [2] S. Zeriuoh, M. Boutahri, S. El Yamani, A. Roukhe, Targets Classification on Multispectral Images using Connectionists Methods, Academia.Edu (n.d.). <https://www.academia.edu/download/82589859/10.11648.j.ajpa.20150303.14.pdf> (accessed April 4, 2025).
- [3] O. bin Laden, ISR offensif dans la guerre contre le terrorisme, Airuniversity.Af.Edu (n.d.). https://www.airuniversity.af.edu/Portals/10/ASPJ_French/journals_O/Volume-02_Issue-1/danskine.pdf (accessed April 4, 2025).
- [4] D. Micle, F. Deiac, A. Olar, R. Drența, C.F.- Agriculture, undefined 2021, Research on innovative business plan. Smart cattle farming using artificial intelligent robotic process automation, Mdpi.Com (n.d.). <https://www.mdpi.com/2077-0472/11/5/430> (accessed April 4, 2025).
- [5] W. Aryatwijuka, H. Mutebi, P. Nagawa, B. Tukamuhabwa, Artificial Intelligence and Humanitarian Supply Chain Resilience: Mediating Effect of Localized Logistics Capacity, (2024). https://www.researchgate.net/profile/Samuel-Mayanja/publication/384941643_Artificial_Intelligence_and_Humanitarian_Supply_Chain_Resilience_Mediating_Effect_of_Localized_Logistics_Capacity/links/67122cc709ba2d0c7606f856/Artificial-Intelligence-and-Humanitarian-Supply-Chain-Resilience-Mediating-Effect-of-Localized-Logistics-Capacity.pdf (accessed April 4, 2025).
- [6] B. Claverie, G.D.-A.J. of Management, undefined 2022, C2-Command and Control: A System of Systems to Control Complexity, Researchgate.Net (n.d.). https://www.researchgate.net/profile/Gilles-Desclaux/publication/363583535_C2_-_Command_and_Control_A_System_of_Systems_to_Control_Complexity/links/639cb203e42faa7e75cb0d61/C2-Command-and-Control-A-System-of-Systems-to-Control-Complexity.pdf (accessed April 4, 2025).
- [7] M.A.-I.T. on automatic control, undefined 2003, Command and control (C2) theory: A challenge to control science, Ieeexplore.Ieee.Org (n.d.). <https://ieeexplore.ieee.org/abstract/document/1104607/> (accessed April 4, 2025).
- [8] A. Hamdouchi, A. Idri, Enhancing IoT security through boosting and feature reduction techniques for multiclass intrusion detection, Neural Computing and Applications 2025 (2025) 1–24. <https://doi.org/10.1007/S00521-025-11001-2>.
- [9] A. Hamdouchi, A. Idri, Comprehensive Evaluation of Federated Learning Configurations for Intrusion Detection in IoT Contexts, 2024 World Conference on Complex Systems (WCCS) (2024) 1–6. <https://doi.org/10.1109/WCCS62745.2024.10765581>.
- [10] N. Kandybko, O. Zemlina, ... V.P.-A. conference, undefined 2023, Application of artificial intelligence technologies in the context of military security and civil infrastructure protection, Pubs.Aip.Org (n.d.). <https://pubs.aip.org/aip/acp/article-abstract/2476/1/030019/2891077> (accessed April 4, 2025).
- [11] G. Sakr, O. Hakme, The Impact of Artificial Intelligence on Military and Security in the MENA, AI in the Middle East for Growth and Business (2025) 363–383. https://doi.org/10.1007/978-3-031-75589-7_21.
- [12] Strategy for the Digital Transformation of UN Peacekeeping | United Nations Peacekeeping, (n.d.).

- <https://peacekeeping.un.org/en/strategy-digital-transformation-of-un-peacekeeping> (accessed April 5, 2025).
- [13] T. Ige, A. Kolade, O. Kolade, Enhancing Border Security and Countering Terrorism Through Computer Vision: A Field of Artificial Intelligence, *Lecture Notes in Networks and Systems* 597 LNNS (2023) 656–666.
https://doi.org/10.1007/978-3-031-21438-7_54.
 - [14] M. Sharma, C.R.S. Kumar, Machine Learning-Based Smart Surveillance and Intrusion Detection System for National Geographic Borders, *Lecture Notes in Electrical Engineering* 806 (2022) 165–176.
https://doi.org/10.1007/978-981-16-6448-9_19.
 - [15] Search and Rescue, (n.d.).
<https://www.frontex.europa.eu/what-we-do/operations/search-and-rescue/> (accessed April 5, 2025).
 - [16] A. Bin Rashid, A.K. Kausik, A. Al Hassan Sunny, M.H. Bappy, Artificial Intelligence in the Military: An Overview of the Capabilities, Applications, and Challenges, *International Journal of Intelligent Systems* 2023 (2023).
<https://doi.org/10.1155/2023/8676366>.
 - [17] M. Hutson, The opacity of artificial intelligence makes it hard to tell when decision-making is biased, *IEEE Spectr* 58 (2021) 40–45.
<https://doi.org/10.1109/MSPEC.2021.9340114>.
 - [18] A. Mensah-Sackey, H. Giezendanner, P. Holtom, Aperçu de la gestion des armes et des munitions en Afrique: rapport sur l'état d'avancement 2023, (2023).
<https://unidir.org/publication/aperçu-de-la-gestion-des-armes-et-des-munitions-en-afrique-rapport-sur-letat-davancement-2023/> (accessed April 5, 2025).
 - [19] J. Souza, R. Avelino, ... S. da S.-A. and B.D. in E., undefined 2023, 11. Artificial intelligence: dependency, Books.Google.Com (n.d.).
<https://books.google.com/books?hl=en&lr=&id=HuDoEAAAQBAJ&oi=fnd&pg=PA228&dq=Artificial+intelligence+Technological+dependence&ots=KmmIbW0VSc&sig=D-COzHMgstFux8ZiN0F6ZC6Nnk> (accessed April 5, 2025).
 - [20] M.W. Khan, M.A. Destek, Z. Khan, Income Inequality and Artificial Intelligence: Globalization and age dependency for developed countries, *Soc Indic Res* (2025).
<https://doi.org/10.1007/S11205-024-03493-7>.
 - [21] H. Baniecki, P.B.-I. Fusion, undefined 2024, Adversarial attacks and defenses in explainable artificial intelligence: A survey, Elsevier (n.d.).
<https://www.sciencedirect.com/science/article/pii/S1566253524000812> (accessed April 5, 2025).
 - [22] S. Qiu, Q. Liu, S. Zhou, C.W.-A. Sciences, undefined 2019, Review of artificial intelligence adversarial attack and defense technologies, Mdpi.Com (n.d.).
<https://www.mdpi.com/2076-3417/9/5/909> (accessed April 5, 2025).
 - [23] Intelligence artificielle au Maroc: l'UNESCO énonce des recommandations | SNRT News, (n.d.).
<https://snrtnews.com/fr/article/intelligence-artificielle-au-maroc-lunesco-enonce-des-recommandations-97072> (accessed April 5, 2025).
 - [24] A. Das, G. Student Member, P. Rad, S. Member, Opportunities and Challenges in Explainable Artificial Intelligence (XAI): A Survey, (2020).
<https://arxiv.org/abs/2006.11371v2> (accessed April 5, 2025).
 - [25] UNESCO - EU Partnership Grows: Tackling Disinformation and Hate Speech Globally | UNESCO, (n.d.).
<https://www.unesco.org/en/articles/unesco-eu-partnership-grows-tackling-disinformation-and-hate-speech-globally> (accessed April 5, 2025).
 - [26] U. Nations, Pacte numérique mondial | Nations Unies, (n.d.).

<https://www.un.org/fr/summit-of-the-future/global-digital-compact> (accessed April 5, 2025).

- [27] Zouinar, Moustafa, Évolutions de l'Intelligence Artificielle : quels enjeux pour l'activité humaine et la relation Humain-Machine au travail ?, <Http://Journals.Openedition.Org/Activites> (2020).
<https://doi.org/10.4000/ACTIVITES.4941>.