

17 Marzo 2026

Queste non sono aziende “hi-tech”, ma contractors militari

di Avner Gvaryahu



Esiste una strategia militare israeliana chiamata “procedura nebbia”. Utilizzata per la prima volta durante la seconda intifada, è una regola non ufficiale che richiede ai soldati di guardia nei posti militari in condizioni di scarsa visibilità di sparare raffiche di proiettili nell’oscurità, basandosi sulla teoria che potrebbe nascondersi una minaccia invisibile.

È violenza autorizzata dalla cecità. Sparare nel buio e chiamarla deterrenza. Con l’alba della guerra basata sull’IA, quella stessa logica della cecità scelta è stata perfezionata, sistematizzata e affidata a una macchina.

La recente guerra di Israele a Gaza è stata descritta come la prima grande “guerra dell’IA” – la prima guerra in cui i sistemi di IA hanno svolto un ruolo centrale nella generazione della lista israeliana dei presunti militanti di Hamas e della Jihad islamica da colpire. Sistemi che elaboravano miliardi di punti dati per classificare la probabilità che una data persona nel territorio fosse un combattente.

L’oscurità nella torre di guardia era una condizione del terreno. L’oscurità dentro l’algoritmo è una condizione della progettazione. In entrambi i casi, la cecità è stata scelta. È stata scelta perché la cecità è utile: crea una possibilità di negazione, fa sembrare la violenza inevitabile, sposta la questione di chi abbia deciso da una persona a una procedura. La nebbia non si è diradata. Le è stato assegnato un punteggio di probabilità e chiamata intelligenza.

Potrebbe essere stata la cecità scelta a portare, all’inizio della guerra USA-Israele contro l’Iran, al bombardamento della scuola elementare Shajareh Tayyebbeh a Minab, nel sud dell’Iran. Almeno 168 persone sono state uccise, la maggior parte delle quali bambini, bambine dai sette ai dodici anni.

Le armi erano precise. Esperti di munizioni hanno descritto il targeting come “incredibilmente accurato”, ogni edificio colpito singolarmente, nulla mancato. Il problema non era l’esecuzione. Il problema era l’intelligence. La scuola era stata separata da una vicina base dei Guardiani della Rivoluzione da una recinzione e riconvertita per uso civile quasi un decennio fa. Da qualche parte nel ciclo di targeting, sembra che questo fatto non sia mai stato aggiornato.

Gaza era il laboratorio. Minab è il mercato

Il ruolo esatto dell'IA nel bombardamento di Minab non è stato ufficialmente confermato. Quello che si sa è che l'infrastruttura di targeting in cui operano questi sistemi non ha alcun meccanismo affidabile per segnalare quando l'intelligence sottostante è vecchia di un decennio.

Che un algoritmo abbia o meno selezionato questa scuola, è stata selezionata da un sistema che il targeting algoritmico ha costruito. Per colpire 1.000 obiettivi nelle prime 24 ore della campagna in Iran, i militari statunitensi hanno fatto affidamento su sistemi di IA per generare, dare priorità e classificare la lista degli obiettivi a una velocità che nessun team umano potrebbe replicare.

Gaza era il laboratorio. Minab è il mercato. Il risultato è un mondo in cui le decisioni di targeting più consequenziali nella guerra moderna sono prese da sistemi che non possono spiegarsi, forniti da aziende che non rispondono a nessuno, in conflitti che non generano né responsabilità né resa dei conti. Questo non è un fallimento del sistema. Questo è il sistema.

Dovremmo resistere alla tentazione di incolpare esclusivamente l'algoritmo per la logica che rende i bambini un tasso di errore accettabile. Nel luglio 2014, quattro ragazzi della famiglia Bakr – Ismail, Zakariya, Ahed e Mohammad, dai nove agli undici anni – furono uccisi su una spiaggia di Gaza. Non era coinvolta alcuna IA.

Il sito era stato pre-classificato come compound navale di Hamas. I ragazzi furono segnalati come sospetti perché correvano, poi camminavano – un comportamento che corrispondeva a un modello di targeting per combattenti che cercano di non attirare l'attenzione. Quando il primo missile colpì, i bambini

sopravvissuti fuggirono. Il drone li seguì e sparò di nuovo. Un ufficiale in seguito testimoniò che da una visuale aerea verticale è molto difficile identificare i bambini. Il bombardamento fu archiviato come errore di targeting.

Un database militare israeliano classificato, esaminato dal *Guardian*, da *+972 Magazine* e da *Local Call*, indicava che degli oltre 53.000 morti registrati a Gaza, i combattenti nominati di Hamas e della Jihad islamica rappresentavano circa il 17%. Ciò suggerisce che il resto, l'83%, fosse costituito da civili. Queste non sono le statistiche di una guerra combattuta con precisione, questa è una guerra in cui l'imprecisione è l'obiettivo. (L'IDF ha contestato le cifre presentate nell'articolo del *Guardian*, ma non ha identificato quali cifre siano "giuste").

Quindi i sistemi di targeting IA non hanno inventato questa logica. L'hanno ereditata, codificata in milioni di punti dati e automatizzata oltre ogni significativo controllo umano. Quando una scuola a Minab viene classificata in un database come compound militare, questo non è un malfunzionamento. È la "procedura nebbia", la stessa logica che ha inseguito quattro ragazzi su una spiaggia di Gaza – che funziona esattamente come progettato, ma a una scala diversa, in un paese diverso, con un'arma diversa. L'oscurità ora ha solo hardware migliore.

Molti di questi sistemi di IA sfidano intrinsecamente il diritto internazionale umanitario, che non richiede semplicemente risultati corretti dalle operazioni militari; richiede un processo attento prima che vengano eseguite.

Un comandante deve compiere ogni ragionevole sforzo per verificare che un obiettivo sia un legittimo obiettivo militare. La

legge richiede anche che si faccia tutto il possibile per proteggere i civili dagli effetti dell'attacco, non come un "ripensamento", ma come un *obbligo* parallelo e paritario.

Tale obbligo non può essere delegato a un sistema il cui ragionamento è opaco e i cui risultati non possono essere interrogati in tempo reale. A Gaza, un algoritmo elaborava i dati su ogni persona nella striscia – registri telefonici, modelli di movimento, connessioni sociali, segnali comportamentali – e produceva una lista classificata di nomi, a ciascuno assegnato un punteggio di probabilità che indicava la probabilità che fosse un combattente.

Questo non è come un analista umano che identifica un noto militante e programma un'arma per colpirlo. L'IA non stava confermando identità. Le stava inferendo, statisticamente, su un'intera popolazione, generando obiettivi che nessun umano aveva valutato individualmente prima che apparissero sulla lista.

La verifica, in questo sistema, significava che un operatore umano rivedeva ogni nome per una media di circa 20 secondi, quanto bastava per confermare che il target fosse maschio. Poi autorizzava. Un solo sistema ha prodotto più di 37.000 obiettivi nelle prime settimane di guerra. Un altro era in grado di generare 100 potenziali siti di bombardamento al giorno. Gli umani nel circuito non stavano esercitando giudizio. Stavamo gestendo una coda.

In Iran, il quadro è, al momento, meno documentato. Ma la scala racconta la sua storia. Due fonti hanno confermato a NBC News che i sistemi di IA di Palantir, che attingono in parte alla tecnologia dei grandi modelli linguistici, sono stati usati per

identificare obiettivi. (L'amministratore delegato di Palantir, Alex Karp, ha detto di "*non poter entrare nei dettagli*" quando gli è stato chiesto di questo su CNBC, ma ha detto che Claude era ancora integrato nei sistemi Palantir usati nella guerra in Iran).

Brad Cooper, capo del Comando Centrale degli Stati Uniti, si è vantato che i militari stanno usando l'IA in Iran per "*setacciare enormi quantità di dati in pochi secondi*" al fine di "*prendere decisioni più intelligenti più velocemente di quanto il nemico possa reagire*". Che ogni bombardamento sia stato o meno assistito dall'IA, il ritmo della campagna è stato possibile solo perché il targeting era stato sostanzialmente automatizzato.

Quando i tempi di verifica riportati per i target assistiti da IA sono misurati in secondi, non stiamo più parlando di giudizio umano con assistenza algoritmica. Stiamo parlando di timbrare l'output di una macchina. E quando i dati di quella macchina sono vecchi di un decennio, le conseguenze sono scritte in file di piccole bare.

Le aziende implicate in questo non sono oscure startup della difesa. Palantir, fondata con i primi finanziamenti della CIA e ora uno dei principali fornitori di infrastrutture IA per i militari statunitensi, ha fornito i sistemi usati nella campagna in Iran. Questi sistemi attingono in parte a Claude di Anthropic, un grande modello linguistico la cui azienda madre ha tentato di resistere alle pressioni del Pentagono per rimuovere i vincoli etici al suo uso per il targeting. Il Pentagono ha risposto minacciando di tagliare i legami e rivolgendosi invece a OpenAI e ad altri. Il mercato per l'uccisione su larga scala non manca di fornitori.

L'episodio è istruttivo: l'unica azienda che ha cercato di tracciare una linea rossa è stata messa da parte e l'uccisione è continuata

senza interruzioni. Google, nonostante significative proteste interne dei dipendenti, ha firmato il Progetto Nimbus, un contratto di cloud computing e IA con il governo e l'esercito israeliano del valore di oltre 1 miliardo di dollari.

Amazon è co-firmataria del Progetto Nimbus insieme a Google. Microsoft aveva una profonda integrazione con i sistemi militari israeliani prima di ritirarsi parzialmente sotto pressione nel 2024, momento in cui i dati sono migrati su Amazon Web Services nel giro di pochi giorni.

Anduril, fondata da Palmer Luckey e con personale composto in gran parte da ex funzionari della difesa statunitense, costruisce sistemi d'arma autonomi esplicitamente progettati per il targeting letale. OpenAI, che fino a poco tempo fa proibiva l'uso militare nei suoi termini di servizio, ha silenziosamente rimosso tale restrizione all'inizio del 2024 e da allora ha cercato contratti con il Pentagono.

Queste sono tra le aziende di maggior valore al mondo, con prodotti di consumo usati da centinaia di milioni di persone, partnership di ricerca universitarie e una significativa influenza politica a Washington, Bruxelles e oltre.

Certo, le aziende private hanno rifornito i militari per secoli – con radio, camion, navigazione satellitare, tecnologia a microonde e, naturalmente, sistemi d'arma complessi. Questo non è nuovo o intrinsecamente corrotto. Il problema del “dual use” è vecchio quanto l'industrializzazione: quasi tutte le tecnologie potenti possono essere usate per scopi militari.

Come far scoppiare la bolla israeliana

Ma il targeting tramite IA non è semplicemente un componente che i militari incorporano nelle loro operazioni. È l'architettura decisionale stessa la cosa che determina chi viene ucciso e perché. Quando un singolo sistema può generare decine di migliaia di obiettivi nel tempo che ci sarebbe voluto a una squadra di intelligence umana per verificarne 10, la domanda non è se le aziende private debbano rifornire i militari. È se qualsiasi quadro giuridico possa sopravvivere al contatto con quel sistema.

Nel diritto internazionale parliamo di quadri di responsabilità: la catena di risposte possibili che va da una decisione di usare la forza letale alla persona che l'ha autorizzata. Un quadro di responsabilità richiede che qualcuno sia identificabile come decisore, che il suo ragionamento sia ricostruibile a posteriori e che si possa dimostrare che gli obblighi procedurali richiesti dalla legge – valutazione della proporzionalità, verifica, precauzione – sono stati seguiti.

Il targeting tramite IA distrugge sistematicamente ciascuna di queste condizioni. L'attribuzione si dissolve attraverso una catena di ingegneri, comandanti, operatori e fornitori aziendali, ognuno dei quali può puntare il dito verso un altro. Il ragionamento scompare in un punteggio di probabilità che nessun avvocato può verificare e nessun tribunale può controinterrogare.

Il processo collassa in un'approvazione di 20 secondi di una raccomandazione della macchina. E le aziende che hanno costruito e venduto il sistema rimangono completamente al di fuori del quadro giuridico, perché il diritto internazionale umanitario è stato progettato per gli stati e i loro agenti, e Palantir non è firmataria delle Convenzioni di Ginevra.

Il quadro di responsabilità non è stato semplicemente messo a dura prova o testato dalla guerra con l'IA. *È stato reso strutturalmente irrilevante.*

Sollevare la nebbia della guerra

Dovremmo smettere di chiamare queste aziende “aziende tecnologiche” e iniziare a chiamarle per quello che sono: *appaltatori della difesa.*

Le più grandi aziende di IA non sono fornitori neutrali di infrastrutture che hanno casualmente trovato un cliente militare. Stanno venendo integrate nell'architettura di targeting della guerra moderna. I loro sistemi sono inseriti nella catena di uccisione, i loro ingegneri hanno il “nulla osta di sicurezza”, i loro dirigenti attraversano la stessa porta girevole che ha sempre collegato la Silicon Valley al Pentagono.

Questi fornitori di IA sono all'avanguardia del complesso militare-industriale e dovrebbero essere regolamentati come tali. Una chiara catena di responsabilità si applica a aziende come Raytheon e Lockheed Martin – comportando controlli sulle esportazioni, supervisione del Congresso, quadri di responsabilità civile e condizioni di approvvigionamento – mentre le deboli normative che si applicano alle aziende che scrivono gli algoritmi che selezionano obiettivi militari non sono mai state applicate, testate o fatte rispettare.

Questa non è una dimenticanza. È *una scelta*, attivamente mantenuta tramite lobbying, tramite la deliberata confusione tra prodotti “commerciali” e “della difesa”, e tramite una cultura normativa che tratta ancora l'IA come una tecnologia di consumo che si è trovata casualmente sul campo di battaglia.

Palantir ha speso quasi 6 milioni di dollari in lobbying a Washington nel 2024, e in un trimestre del 2023 ha superato in spesa Northrop Grumman. Ha lanciato una fondazione dedicata per modellare l'ambiente politico in cui opera.

Il consorzio di Palantir, Anduril, OpenAI, SpaceX e Scale AI è stato descritto dai suoi stessi partecipanti come un progetto per fornire una nuova generazione di appaltatori della difesa al governo statunitense. Le società di venture capital che sostengono queste aziende, Andreessen Horowitz e Founders Fund, hanno coltivato influenza attraverso la prossimità al potere: ex alti funzionari nei loro consigli consultivi, soci che ruotano attraverso ruoli governativi e accesso diretto ai decisori politici che determinano quanto il Pentagono spende e per cosa.

L'AI Act dell'UE, il tentativo più ambizioso finora di governare l'intelligenza artificiale, esenta esplicitamente le applicazioni militari e di sicurezza nazionale, con la giustificazione dichiarata che il diritto internazionale umanitario è il quadro più appropriato. È un notevole atto di circolarità: *l'unico corpo giuridico che viene sistematicamente distrutto da questi sistemi è designato come loro regolatore, mentre i regolatori che potrebbero effettivamente limitarli distolgono lo sguardo.*

Negli Stati Uniti, le disposizioni sull'IA del National Defense Authorization Act 2025 non regolano l'IA militare. Dirigono le agenzie ad adottarne di più. La strategia IA di Pete Hegseth, emanata nel gennaio 2026, inquadra la questione interamente come una gara, dirigendo il Pentagono a muoversi a velocità di guerra, con l'IA come primo banco di prova. La cultura normativa non è semplicemente rimasta indietro rispetto alla tecnologia. Ha deciso, deliberatamente, di non provarci.

Finora, l'unico serio intervento governativo sulla capacità militare dell'IA che abbiamo visto non è venuto da uno Stato che chiedeva moderazione o responsabilità, ma dagli *Stati Uniti che chiedevano che i sistemi fossero resi più letali*. Questo è l'orizzonte di ambizione che abbiamo accettato.

Mettere al bando questi sistemi è impossibile quando così tanti degli attori coinvolti si preoccupano poco del diritto internazionale. Ma i punti di pressione rimangono, e sono reali. Qualsiasi futuro governo a Washington che voglia usare la capacità militare dell'IA senza produrre una serie infinita di Minab avrà bisogno di un quadro normativo – non come concessione ai critici, ma come requisito di base per non diventare un “attore canaglia”.

Lo stesso vale in Europa, dove la Gran Bretagna ha impegnato oltre 1 miliardo di sterline in un nuovo sistema di targeting integrato con l'IA che collega sensori e capacità d'attacco in tutti i domini, e dove la principale azienda IA francese ha collaborato con una startup tedesca della difesa per costruire piattaforme d'arma autonome, e dove la Germania sta schierando droni d'attacco guidati dall'IA in Ucraina.

C'è un'opportunità per regolamentare questi sistemi. L'UE ha gli strumenti più ovvi, non attraverso l'AI Act, che deliberatamente esenta le applicazioni militari, ma attraverso i controlli sulle esportazioni e le condizioni di approvvigionamento per i sistemi a duplice uso che si muovono tra i mercati commerciali e della difesa.

Anche le corti internazionali stanno iniziando ad aprire porte: il parere consultivo della Corte Internazionale di Giustizia sui diritti

dei palestinesi ha creato un quadro in cui le aziende che forniscono sistemi usati in attacchi illegittimi affrontano una potenziale esposizione alla responsabilità in giurisdizioni che prendono sul serio il diritto internazionale. E le aziende di IA hanno bisogno dei governi, non solo come clienti, ma come fornitori della potenza di calcolo, dell'energia e dell'infrastruttura fisica che l'IA all'avanguardia richiede e che nessuna azienda può sostenere con le sole entrate commerciali.

Questa dipendenza dà agli Stati che sono disposti a usarla una reale influenza sulle aziende che preferirebbero non essere regolamentate. La domanda è se un governo con gli strumenti per agire deciderà, prima della prossima Minab, che il costo dell'inazione è diventato troppo alto.

L'aspetto che la regolamentazione dovrebbe avere è relativamente semplice, anche se è difficile da applicare. I sistemi di IA usati nel targeting devono essere spiegabili – non tramite un punteggio di probabilità, ma con un ragionamento che un avvocato possa verificare.

Il costo civile cumulativo delle campagne assistite dall'IA deve essere valutato nel suo insieme. E la responsabilità che si ferma all'operatore deve estendersi lungo la catena di fornitura fino alle aziende che hanno consapevolmente costruito e venduto sistemi opachi per l'uso in conflitti armati. Queste non sono richieste nuove. Sono le condizioni minime affinché le leggi di guerra significhino qualcosa nell'era del targeting algoritmico.

Nel frattempo, la “procedura nebbia” è operativa e sta definendo il futuro della guerra. Ma i soldati che sparavano nell'oscurità, almeno, erano presenti in essa. Le aziende che hanno costruito ciò

che li ha sostituiti lo fanno da Palo Alto, senza rischi personali, senza esposizione legale e con ogni incentivo a rifarlo.

* da *The Guardian* – Avner Gvaryahu è un ricercatore di dottorato presso la Blavatnik School of Government dell'Università di Oxford. È stato direttore esecutivo di Breaking the Silence, un'organizzazione israeliana per i diritti umani composta da ex soldati.