

UNIVERSIDADE FEDERAL DE UBERLÂNDIA
INSTITUTO DE ECONOMIA E RELAÇÕES INTERNACIONAIS
CURSO DE RELAÇÕES INTERNACIONAIS

HENRIQUE CAIRES RIBEIRO

INTELIGÊNCIA ARTIFICIAL E TOMADA DE DECISÃO:
o caso dos sistemas Lavender e Gospel no conflito entre Israel e Hamas (2023–2024)

UBERLÂNDIA

2026

HENRIQUE CAIRES RIBEIRO

INTELIGÊNCIA ARTIFICIAL E TOMADA DE DECISÃO:

o caso dos sistemas Lavender e Gospel no conflito entre Israel e Hamas (2023–2024)

Trabalho de Conclusão de Curso apresentado ao curso de graduação em Relações Internacionais do Instituto de Economia e Relações Internacionais da Universidade Federal de Uberlândia, na categoria de Projeto de Mestrado, como requisito para a obtenção do título de bacharel em Relações Internacionais.

Orientador: Prof. Dr. Edson José Neves Júnior.

Área de concentração: Segurança Internacional / Estudos Estratégicos / Tecnologias Emergentes.

UBERLÂNDIA

2026

HENRIQUE CAIRES RIBEIRO

INTELIGÊNCIA ARTIFICIAL E TOMADA DE DECISÃO:

o caso dos sistemas Lavender e Gospel no conflito entre Israel e Hamas (2023–2024)

Trabalho de Conclusão de Curso apresentado ao curso de graduação em Relações Internacionais do Instituto de Economia e Relações Internacionais da Universidade Federal de Uberlândia, na categoria de Projeto de Mestrado, como requisito para a obtenção do título de bacharel em Relações Internacionais.

Uberlândia, 27 de julho de 2026.

Banca examinadora:

Prof. Dr. Edson José Neves Júnior (Orientador) – UFU

Prof. Dr. Flávio Pedroso Mendes - UFU

Prof. Dr. João Fernando Finazzi - UFU

RESUMO

A integração de sistemas de inteligência artificial às funções militares, a partir de sua essência como tecnologia de uso dual, desloca o eixo de transformação da guerra contemporânea da plataforma de armas para o processo decisório que antecede o emprego da força. Este projeto tem como objetivo analisar em que medida, e sob que condições, o emprego dos sistemas Lavender e Gospel pelas Forças de Defesa de Israel (IDF), no conflito contra o Hamas (2023–2024), afetou o escrutínio humano sobre a decisão de uso da força letal. O objeto da pesquisa é, assim, examinar o grau de deliberação humana exercido nessa decisão, em sua relação com sistemas algorítmicos de apoio à seleção de alvos. Ancorada na perspectiva dos Estudos Estratégicos — e deliberadamente afastada de enfoques normativos — a pesquisa parte da hipótese de que o emprego desses sistemas teria promovido uma transferência parcial da autoridade decisória do operador humano para o sistema algorítmico, com redução do grau de deliberação humana efetiva sobre a seleção de alvos. O estudo estrutura-se como pesquisa qualitativa de caso único, baseada em documentação secundária de fontes abertas submetidas a verificação cruzada, e desenvolve-se procedimentalmente em etapas que vão da constituição do corpus documental à avaliação da hipótese segundo critérios de decisão explícitos. A análise confronta as evidências do caso com um quadro conceitual estruturado em torno das noções de controle humano significativo, viés de automação, cadeia de letalidade, seletor e tomada de decisão algorítmica, examinando o controle humano em suas dimensões temporal, informacional e de autoridade efetiva. Espera-se, com isso, contribuir para o debate sobre os limites do controle humano em sistemas contemporâneos de apoio à decisão militar, oferecendo evidências empíricas sobre a reconfiguração da autoridade decisória no emprego algorítmicamente mediado da força letal.

Palavras-chave: inteligência artificial; tecnologia de uso dual; controle humano significativo; Forças de Defesa de Israel; Hamas.

ABSTRACT

The integration of artificial intelligence systems into military functions, stemming from their essence as a dual-use technology, shifts the axis of transformation of contemporary warfare from the weapons platform to the decision-making process that precedes the use of force. This research project aims to analyse to what extent, and under which conditions, the employment of the Lavender and Gospel systems by the Israel Defense Forces (IDF), in the conflict against Hamas (2023–2024), affected human scrutiny over the decision on the use of lethal force. The object of the research is therefore to examine the degree of human deliberation exercised in that decision, in its relation to algorithmic systems of target-selection support. Anchored in the perspective of Strategic Studies — and deliberately detached from normative approaches — the research departs from the hypothesis that the employment of these systems would have promoted a partial transfer of decision-making authority from the human operator to the algorithmic system, reducing the degree of effective human deliberation over target selection. The study is designed as qualitative single-case research, based on secondary documentation from open sources submitted to cross-verification, and unfolds procedurally in stages ranging from the constitution of the documentary corpus to the assessment of the hypothesis according to explicit decision criteria. The analysis confronts the evidence of the case with a conceptual framework structured around the notions of meaningful human control, automation bias, kill chain, selector and algorithmic decision-making, examining human control in its temporal, informational and effective-authority dimensions. It is expected to contribute to the debate on the limits of human control in contemporary military decision-support systems, offering empirical evidence on the reconfiguration of decision-making authority in the algorithmically mediated employment of lethal force.

Keywords: artificial intelligence; dual-use technology; meaningful human control; Israel Defense Forces; Hamas.

1 INTRODUÇÃO

O conflito armado contemporâneo atravessa uma inflexão cujo eixo não é mais apenas a plataforma de armas, mas o processo decisório que a antecede. Ao longo do século XX, as grandes transformações militares — como a Revolução em Assuntos Militares (RMA) do final do século — foram lidas predominantemente pela ótica da capacidade material: alcance, precisão, letalidade e mobilidade dos sistemas de armas. O reconhecimento da inteligência artificial (IA) como tecnologia de uso dual — isto é, uma tecnologia que opera tanto na utilidade civil como na militar — desloca essa leitura: o que está em transformação vai além de como a força é aplicada e alcança quem, ou o quê, participa da decisão de aplicá-la.

Assim, sistemas algorítmicos de apoio à decisão passam a operar nas etapas que antecedem o disparo — coleta e análise de dados, identificação e classificação de alvos, recomendação de engajamento — comprimindo os intervalos do ciclo OODA tradicional (Observar, Orientar, Decidir e Agir) (Scharre, 2018). Essa compressão temporal e a delegação progressiva de tarefas cognitivas a máquinas constituem o problema estratégico central que este projeto se propõe a investigar.

A lógica de observação aplicada aqui é orientada a partir dos Estudos Estratégicos. Isso significa tratar a IA militar não como artefato técnico isolado nem como objeto de avaliação normativa, ou seja, passa a ser entendida como uma variável que incide sobre a dinâmica de poder, sobre a cadeia de comando e sobre a relação entre decisão humana e produção de violência letal. Vale ressaltar que a distinção é deliberada, de forma em que, reconhece-se que parte da literatura recente se dedica à discussão da IA militar pela via da conformidade jurídica ou da regulação internacional; a este trabalho, contudo, interessa compreender de que maneira a incorporação dessas tecnologias reconfigura o controle humano efetivo sobre a decisão de uso da força — um problema que pertence ao núcleo dos Estudos Estratégicos porque toca diretamente a estrutura de autoridade, a velocidade do ciclo decisório e a distribuição da capacidade de coerção.

O caso que materializa esse problema é o emprego dos sistemas Lavender e Gospel pelas Forças de Defesa de Israel (IDF) no conflito contra o Hamas, no período de 2023 a 2024. Cabe explicitar, desde já, que se trata de um conflito de natureza assimétrica, e não de uma guerra convencional interestatal, de modo em que, no caso, tem-se de um lado, um ator estatal com força aérea e domínio tecnológico; de outro, um ator não estatal. Essa característica não fragiliza o recorte, mas o delimita, ao esclarecer que o objeto de interesse se restringe à IA na função específica de geração e seleção de alvos — e não ao seu emprego na guerra em abstrato — fenômeno que se manifesta com particular densidade documental justamente no caso em

questão. Nesse sentido, trata-se de um caso com a materialidade documental necessária, uma vez que a utilização desses sistemas foi objeto de investigações jornalísticas detalhadas, de relatórios e de análise acadêmica, produzindo um conjunto de fontes abertas que permite caracterizar, ainda que parcialmente, o impacto desses sistemas sobre a posição do operador humano na cadeia de uso da força.

Dessa forma, nesse conjunto de fontes, Faris (2025) quantifica operações que empregaram o sistema Lavender no período: a tecnologia teria processado dados de aproximadamente 2,3 milhões de habitantes de Gaza, atribuindo a cada indivíduo uma pontuação de probabilidade de vínculo com organizações armadas e gerando, em períodos de pico, cerca de 37 mil alvos potenciais. A literatura sobre o caso indica, ainda, que a verificação humana de cada alvo consumia, em média, cerca de vinte segundos — tempo destinado, sobretudo, a confirmar o gênero do indivíduo, e não a escrutinar a lógica de seleção ou a avaliar fatores de risco (Abraham, 2024; Faris, 2025). Esse arranjo é corroborado por análises acadêmicas independentes (Downey, 2025; Reichert, 2025), que reportam os mesmos parâmetros de escala e de verificação. Em conjunto, tais observações descrevem uma lógica na qual a escala e a velocidade da produção de alvos passam a condicionar a própria possibilidade da supervisão humana que deveria avaliar danos colaterais e proporcionalidade — o que, por si, contrasta com o discurso de maior precisão associado à mera incorporação de tecnologia ao conflito armado contemporâneo.

Com isso, esta pesquisa tem como objeto o grau de deliberação humana exercido sobre a decisão de uso da força letal, em sua relação com o emprego de sistemas algorítmicos de apoio à seleção de alvos. A formulação merece ser lida em seus dois polos. De um lado, há um fato preexistente e incontroverso: existe uma decisão sobre o uso da força letal, e essa decisão era, até recentemente, exercida sob controle humano direto. De outro, há um elemento novo que se interpõe nesse processo: sistemas algorítmicos que passam a informar, acelerar e estruturar essa decisão. Com isso, o que se investiga é, precisamente, a relação entre esses dois polos — quanto, e de que maneira, a presença do sistema algorítmico altera o grau de controle que o operador humano efetivamente exerce — e, com isso, busca-se qualificar o peso dessa participação sobre a margem de escrutínio humano.

Convém antecipar, por fim, a natureza precisa do problema, de forma em que, se situa aquém da automação plena do engajamento — a ideia da “máquina que mata sozinha” — numa zona intermediária em que o operador humano permanece formalmente no circuito decisório, porém sob condições — de tempo, de informação e de estrutura cognitiva — que podem tornar

seu controle progressivamente menos significativo. É exatamente essa zona intermediária que o caso Lavender/Gospel materializa e que o trabalho pretende examinar.

2 JUSTIFICATIVA

A relevância deste estudo decorre de três ordens de consideração: a magnitude da transformação em curso, a posição estratégica que a tecnologia em questão ocupa na disputa internacional por poder militar e a tentativa de mitigar uma lacuna de produção acadêmica específica acerca do tema.

Em primeiro lugar, a incorporação de inteligência artificial às funções militares deixou de ser projeção — ou exceção — para se tornar variável recorrente do conflito armado contemporâneo. Nesse cenário, a IA já está integrada a um amplo conjunto de capacidades militares, com impacto direto sobre a utilidade, a acurácia, a letalidade e a autonomia dos sistemas de armas (Erskine; Miller, 2024).

Ainda nesse sentido, é fundamental estabelecer que, mais relevante para este projeto do que o emprego da IA na execução da força — os chamados sistemas de armas autônomas letais (LAWS) — é seu emprego nas etapas anteriores ao disparo, ou seja, na produção de inteligência, na seleção de alvos e na recomendação de engajamento. Sendo assim, é justamente nesse domínio que o conflito recente entre Israel e Hamas oferece um escopo analítico válido, ao caracterizar a passagem de um modelo em que o analista humano constrói o alvo para um cenário no qual um sistema algorítmico baseado em IA gera o alvo e o submete ao operador para mera validação ritualística — o que a literatura acadêmica, a exemplo de Dorsey (2026), caracteriza como fenômeno de *rubber-stamping*¹.

Portanto, o caso investigado apresenta uma mudança qualitativa na arquitetura decisória, que ultrapassa o mero ganho incremental de eficiência. De tal forma, compreender as consequências dessa mudança impõe-se ao campo dos Estudos Estratégicos e da Segurança Internacional, porque as decisões tomadas a partir da incorporação dessas tecnologias podem determinar a estrutura do controle humano no engajamento em combate no futuro.

É preciso, ainda, delimitar com exatidão a natureza dessa função, evitando uma imprecisão que reduziria o caso a mera coleta de inteligência. A geração algorítmica de alvos apoia-se, de fato, em capacidades de vigilância e tratamento de sinais, o que a aproxima, na origem, das atividades de inteligência. O que confere ao caso caráter operacional — e não meramente informacional — é a supressão do intervalo entre a produção do alvo e o seu

¹Dorsey (2026) utiliza o termo *rubber-stamping* — “carimbar” — para referir-se ao fenômeno da validação, sem análise crítica, da recomendação do algoritmo para uso da força letal, o que, segundo a autora, pode acarretar equívocos ou escalada não intencional do combate.

engajamento. Com efeito, na atividade de inteligência clássica, o produto informa uma decisão posterior, mediada e deliberada; já no arranjo aqui examinado, a lista gerada pelo sistema converte-se diretamente em ordem de ataque, com verificação humana da ordem de vinte segundos, sendo tratada, segundo Abraham (2024) e Faris (2025), como se fosse uma decisão humana.

Assim sendo, essa fusão entre a função de inteligência e a função de fogo — a integração entre sensor e atirador que Rangel e Barbosa (2026) designam como vigilância letal — desloca o fenômeno do registro da espionagem para o da condução operacional da força. O objeto do projeto localiza-se, portanto, precisamente nessa fronteira: o ponto em que a produção algorítmica de alvos deixa de ser insumo de análise e passa a estruturar, em tempo comprimido, a decisão de uso da força letal.

Em segundo lugar, a IA militar tornou-se elemento central da competição entre potências: a literatura descreve o estágio atual como uma corrida armamentista global pela inteligência artificial, na qual a necessidade percebida de igualar as capacidades de potenciais adversários funciona como motor de aceleração tecnológica (Erskine; Miller, 2024). Seguindo essa lógica, tal dinâmica é estrutural, e não conjuntural, uma vez que mesmo Estados com forças convencionais enfraquecidas mantêm a IA como prioridade justamente porque ela promete um caminho não linear de redução de assimetrias de poder — lógica observável tanto no investimento das grandes potências quanto na aposta de atores que buscam compensar inferioridade material por meio de capacidades algorítmicas (Zysk, 2023). Nesse contexto, o fenômeno deixa de ser questão doméstica de doutrina e passa a ser variável da competição estratégica entre atores: a compressão dos ciclos decisórios, a delegação de tarefas cognitivas a máquinas e a velocidade da produção algorítmica de alvos têm o potencial de afetar a percepção de intenções, a calibragem da dissuasão e o risco de escalada não intencional (Johnson, 2021). Investigar como esses sistemas reconfiguram o controle humano é, portanto, averiguar um dos vetores que redefinem a distribuição contemporânea do poder militar.

Em terceiro lugar — e é aqui que se localiza a contribuição específica deste projeto — existe uma lacuna no enquadramento dominante do problema. No contexto acadêmico, a maior parte da produção que trata de Lavender e Gospel o faz pela via da denúncia humanitária ou da avaliação de conformidade com normas internacionais (Downey, 2025; Reichert, 2025). Esse enfoque é legítimo, mas deixa em aberto uma questão de que, independentemente do juízo normativo sobre essas operações, as evidências disponíveis sugerem a necessidade de averiguar também o grau efetivo de controle humano exercido na cadeia de seleção de alvos.

Deslocar a pergunta do registro normativo para o registro estratégico permite tratar o caso como aquilo que ele empiricamente é — a operacionalização concreta de um conceito até então majoritariamente teórico, o de controle humano significativo — e examinar, com base em fontes abertas, em que pontos da cadeia decisória o humano permaneceu, em que pontos foi deslocado e como a escala e a velocidade do sistema condicionaram essa transformação. É essa lacuna que o projeto pretende mitigar, oferecendo uma análise empiricamente ancorada sobre o que o caso revela a respeito dos limites do controle humano em sistemas contemporâneos de apoio à decisão militar.

Justifica-se, por fim, a escolha do caso único, pois o emprego de Lavender e Gospel reúne uma combinação de fatores que o torna analiticamente denso: possui materialidade documental — foi reportado, investigado e parcialmente descrito em fontes acessíveis — ocorre em recorte temporal delimitado (2023–2024) e expõe a tensão entre o produto do algoritmo de determinação de alvos e a supervisão humana. Encontra-se, assim, uma oportunidade para observar, em um exemplo concreto, a operacionalização de conceitos que, sem ancoragem empírica, permaneceriam no plano da especulação teórica.

3 PROBLEMÁTICA E HIPÓTESE

A construção do problema parte de uma constatação e de uma transformação. A constatação é a seguinte: existe, na condução da guerra, uma decisão sobre o emprego da força letal — fato que precede qualquer tecnologia e que, até período recente, era exercido sob controle humano direto. A transformação, por sua vez, é a interposição, nessa decisão, de sistemas de inteligência artificial que passam a produzir, classificar e recomendar alvos em escala e velocidade sem precedentes. Com isso, o problema de pesquisa nasce do encontro entre esses dois elementos e pode ser enunciado nos seguintes termos: em que medida, e sob que condições, o emprego dos sistemas Lavender e Gospel afetou o grau de deliberação humana sobre a decisão de uso da força letal no conflito entre Israel e Hamas (2023–2024)?

Nesse sentido, a formulação é relacional, pois não formula juízo de valor — se a inteligência artificial é boa ou má, legítima ou ilegítima — limitando-se a indagar pela natureza e pelo grau de uma eventual reconfiguração. Trata-se de uma relação entre duas variáveis, de um lado, o emprego dos sistemas algorítmicos de apoio à seleção de alvos; de outro, o controle humano efetivamente exercido sobre a decisão de uso da força letal. O problema, assim construído, é cientificamente operável porque suas variáveis são identificáveis e sua relação é, ao menos parcialmente, observável a partir das evidências disponíveis em fontes abertas sobre o caso.

À pergunta corresponde uma resposta provisória, a ser confrontada com as evidências ao longo da pesquisa: o emprego dos sistemas Lavender e Gospel teria promovido uma transferência parcial da autoridade decisória do operador humano para o sistema algorítmico, com redução do grau de deliberação humana efetiva sobre a seleção de alvos.

Cabe explicitar o estatuto desta hipótese: trata-se de resposta inicial, e não de conclusão *a priori*. Dessa forma, a pesquisa se considerará cumprida tanto se as evidências a confirmarem quanto se a refutarem — caso a investigação demonstre que o grau de escrutínio humano permaneceu robusto e efetivo apesar da presença dos sistemas, o trabalho terá igualmente respondido ao seu problema. Portanto, o compromisso da pesquisa é com a pergunta, não com a confirmação da resposta provisória. Os critérios objetivos pelos quais a hipótese será sustentada, qualificada ou refutada são definidos na seção metodológica.

4 OBJETIVOS

4.1 Objetivo geral

Analisar em que medida, e sob que condições, o emprego dos sistemas Lavender e Gospel, pelas Forças de Defesa de Israel, afetou a deliberação humana sobre a decisão de uso da força letal no conflito entre Israel e Hamas (2023–2024).

4.2 Objetivos específicos

São objetivos específicos:

- a) descrever o funcionamento dos sistemas Lavender e Gospel a partir das evidências disponíveis em fontes abertas;
- b) identificar em quais etapas da cadeia de seleção de alvos ocorreu participação humana;
- c) analisar em que medida a velocidade e a escala da produção algorítmica de alvos influenciaram a supervisão humana do processo decisório;
- d) discutir as implicações estratégicas decorrentes da utilização de sistemas algorítmicos no emprego da força letal.

5 FUNDAMENTAÇÃO TEÓRICO-CONCEITUAL

Esta seção delimita o aparato teórico-conceitual que sustenta a investigação e constrói, com precisão, os conceitos que serão mobilizados como instrumentos de análise do caso. Cinco noções organizam a fundamentação: (1) controle humano significativo, (2) viés de automação, (3) cadeia de letalidade (*kill chain*), (4) seletor e (5) tomada de decisão algorítmica. Tomados em conjunto, esses conceitos permitem traduzir a pergunta da pesquisa — em que medida, e sob que condições, a IA afetou o grau de deliberação humana sobre a decisão de uso da força letal — em categorias observáveis. Antes deles, porém, descreve-se o funcionamento

operacional dos dois sistemas, condição para que as afirmações sobre autonomia e controle sejam verificáveis.

A fim de precisar o funcionamento dos sistemas, delimitam-se três aspectos: (i) as bases de dados sobre as quais os sistemas operam, (ii) a posição que ocupam na cadeia de comando e (iii) a natureza da função que exercem — de apoio à decisão, e não de arma autônoma. Tais esclarecimentos são condição para a análise, por tornarem verificáveis as afirmações sobre autonomia e controle que o projeto examina.

Lavender e Gospel são dois sistemas distintos de geração de alvos que operam sobre a mesma base de vigilância. O Gospel produz alvos materiais, ao marcar edifícios e estruturas associados à atuação de operativos. O Lavender, por sua vez, produz alvos humanos, a partir da classificação de indivíduos como supostos operativos do Hamas ou da Jihad Islâmica Palestina, incorporando-os a uma lista de alvos (Abraham, 2023, 2024). Ambos cumprem a mesma função na arquitetura decisória — automatizar a etapa de determinação de alvos — distinguindo-se pelo objeto que geram.

Em relação à base de dados, o Lavender opera sobre a informação de um sistema de vigilância de massa que cobre a quase totalidade dos cerca de 2,3 milhões de habitantes de Gaza (Abraham, 2024; Reichert, 2025). Seu funcionamento segue a lógica do aprendizado de máquina supervisionado: o sistema é treinado com dados de operativos conhecidos, extrai as características (*features*) estatisticamente associadas a esse grupo e as busca na população geral, atribuindo a cada pessoa uma pontuação de 1 a 100, correspondente à probabilidade estimada de vínculo com organização armada. Entre as *features* que elevam a pontuação estão sinais derivados de metadados e de redes sociais — pertencer a um grupo de mensagens com um militante conhecido, trocar de aparelho celular com frequência ou mudar de endereço reiteradamente (Abraham, 2024). Um limiar de pontuação define a conversão em alvo; rebaixá-lo amplia o número de alvos, que atingiu cerca de 37 mil no pico (Abraham, 2024). Fontes acadêmicas registram, ainda, que a unidade responsável declarou uma acurácia de 90% para o sistema, sem que o critério de aferição desse índice tenha sido tornado público (Reichert, 2025). Quanto ao Gospel, as fontes descrevem sua função — a geração acelerada de alvos de estrutura, em ritmo que passou a superar a capacidade de ataque disponível (Abraham, 2023) — em grau de detalhe inferior ao documentado para o Lavender; essa assimetria documental é assumida como limite e será explicitada sempre que a análise depender de dados disponíveis para apenas um dos sistemas.

Na cadeia de comando, os sistemas ocupam a etapa de produção de inteligência e determinação de alvos — as fases de observação e orientação do ciclo OODA — a montante

da decisão de engajamento. A operação encadeia-se em três momentos: primeiro, o Lavender gera o alvo e define o “quem”; em seguida, um sistema de rastreamento associado, o “Where’s Daddy?”, monitora os indivíduos marcados e dispara um alerta quando o alvo entra em sua residência, fixando o “quando” e o “onde” do ataque; por fim, a decisão de atacar e a seleção da munição competem a operadores humanos, situados a jusante da produção algorítmica (Abraham, 2024). O produto do algoritmo alimenta, portanto, as etapas subsequentes e condiciona o disparo.

Daí decorre o terceiro aspecto, decisivo para o enquadramento teórico: Lavender e Gospel não são armas autônomas nem controlam sistema de armas algum. Nesse sentido, são sistemas de apoio à decisão — produzem listas e recomendações, não executam o ataque. O emprego da força cabe a aeronaves tripuladas, e a escolha da munição é decisão humana — segundo as fontes, alvos de baixa patente marcados pelo Lavender foram atacados preferencialmente com munições não guiadas, as chamadas bombas “burras”, de maior efeito destrutivo (Abraham, 2024). Vale registrar, ainda, um dado convergente sobre a interface: relatos indicam que o sistema de rastreamento emite uma recomendação de engajamento que só pode ser abortada por ação manual do operador, sendo o ataque executado por um comando afirmativo — arranjo que estrutura a decisão em favor da aceitação, e não da contestação (Reichert, 2025). Dessa forma, situar corretamente essa função afasta o caso do debate sobre a “máquina que mata sozinha” e o coloca justamente na zona que interessa ao projeto, a de sistemas que estruturam a decisão de uso da força letal a montante, enquanto o acionamento permanece, nominalmente, humano.

No que respeita à arquitetura técnica, as fontes abertas permitem reconstruir a lógica de funcionamento, ainda que não a totalidade dos parâmetros internos. O procedimento corresponde ao de um classificador estatístico treinado a partir de um conjunto rotulado de operativos conhecidos, do qual o sistema deriva um vetor de características e, com base nele, ordena a população por proximidade a esse perfil (Abraham, 2024). Portanto, o resultado não é uma decisão binária, mas uma ordenação probabilística sobre a qual se aplica um limiar operacional ajustável — o que explica por que o rebaixamento do limiar amplia diretamente o volume de alvos (Downey, 2025).

Convém explicitar, contudo, o limite documental: as fontes não detalham a arquitetura interna do modelo, os pesos atribuídos a cada característica nem a composição exata dos dados de treinamento, elementos que permanecem classificados. A pesquisa, portanto, não reconstrói o sistema em nível de engenharia, mas caracteriza sua lógica operacional no grau de precisão

que as evidências sustentam — suficiente para a análise do controle humano, que incide sobre o modo de operação e sobre a interface decisória, e não sobre o código.

Com os aspectos de funcionamento dos sistemas devidamente esclarecidos, faz-se necessário, também, explicitar os conceitos estruturantes da análise do projeto.

Dessa forma, o conceito de controle humano significativo constitui o eixo em torno do qual a investigação se organiza, por ser a variável que se pretende observar em sua relação com a presença do sistema algorítmico. Nesse sentido, a noção emergiu no debate sobre sistemas de armas autônomas para designar uma exigência aparentemente simples, mas analiticamente difusa: a de que as decisões sobre o uso da força letal permaneçam sob controle humano efetivo, e não meramente nominal. A dificuldade conceitual reside precisamente no termo “significativo”, de modo que não basta que um operador humano esteja formalmente presente na cadeia decisória — a presença pode ser uma formalidade vazia se esse operador não dispõe de tempo, informação ou autonomia para exercer julgamento independente sobre a recomendação que recebe (Dorsey, 2026).

A formulação de origem do problema está em Scharre (2018), cuja sistematização da relação humano-máquina ancora analiticamente a exigência de controle. O autor demonstra que a noção de “autônomo” é vazia quando desvinculada da tarefa específica que se automatiza, e organiza o grau de autonomia de qualquer sistema segundo a posição que o operador ocupa diante do ciclo decisório — no caso militar, o ciclo OODA (Scharre, 2018).

Dessa sistematização deriva a tríade que estrutura o debate sobre supervisão humana e delimita as posições possíveis do operador na cadeia. Na operação semiautônoma, o humano permanece *in the loop*, ou seja, a máquina percebe o ambiente e recomenda um curso de ação, mas não executa sem aprovação humana, de modo que o ciclo é interrompido pela decisão do operador. Já na operação supervisionada, o humano situa-se *on the loop*: uma vez posto em funcionamento, o sistema percebe, decide e age por conta própria, cabendo ao operador apenas observar e intervir para interrompê-lo, se assim desejar. Por último, na operação plenamente autônoma, o humano encontra-se *out of the loop*, formato no qual o sistema observa, decide e age inteiramente sem intervenção (Scharre, 2018). É justamente essa gradação que permite situar com precisão, na análise do caso, a posição em que o operador efetivamente se encontrava — dado que a presença formal de um humano *in the loop* não garante, por si, que o controle exercido seja significativo, pois a posição nominal no circuito e a qualidade do escrutínio são coisas distintas.

Para os fins desta pesquisa, decompõe-se o controle humano significativo em três dimensões operacionais, que servirão de critério na análise do caso. A primeira é a dimensão

temporal, segundo a qual o controle só é significativo se o operador dispõe de tempo suficiente para deliberar sobre a recomendação da máquina. A segunda é a dimensão informacional, segundo a qual o controle pressupõe que o operador compreenda a base sobre a qual a recomendação foi gerada e suas limitações. A terceira é a dimensão de autoridade efetiva, segundo a qual o controle exige que o operador possa, de fato, rejeitar ou modificar a recomendação, e não apenas ratificá-la — de forma semelhante ao que Dorsey (2026) caracteriza como *rubber-stamping*. Essa decomposição permitirá, na fase analítica, transformar um conceito abstrato em critérios aferíveis a partir das evidências disponíveis sobre Lavender e Gospel — por exemplo, confrontando a exigência de deliberação temporal com o dado de que a verificação de cada alvo consumia cerca de vinte segundos.

Se o controle humano significativo é a variável a observar, o viés de automação — frequentemente encontrado na literatura como *automation bias* — é o principal mecanismo psicológico-cognitivo que explica sua erosão mesmo quando o humano permanece formalmente presente. Nessa lógica, o conceito designa, na formulação de Johnson (2021), a tendência humana de empregar a automação — os auxílios automatizados de apoio à decisão — como substituto heurístico da busca vigilante de informação, da verificação cruzada e da supervisão adequada do processamento. Vale destacar, ainda, que se trata de um viés documentado na literatura empírica sobre interação humano-máquina, anterior à atual geração de IA militar, observado em contextos que vão da aviação comercial ao apoio à decisão médica (Johnson, 2021).

A relevância do *automation bias* para este projeto está em que ele inverte a presunção — corrente no debate sobre armas autônomas — de que manter um humano no circuito funciona, por si, como salvaguarda. Conforme observam Erskine e Miller (2024), estudos empíricos demonstram que a confiança excessiva em sistemas automatizados torna os decisores humanos menos propensos a mobilizar a própria *expertise* para testar as recomendações da máquina. Os autores destacam, além disso, consequências particularmente relevantes para o campo estratégico: a aceitação do erro, a desqualificação progressiva do julgamento humano (*de-skilling*) e a responsabilidade deslocada — a percepção equivocada de que a máquina pode portar responsabilidade por decisões que são, necessariamente, humanas (Erskine; Miller, 2024).

Assim, o *automation bias* conecta-se diretamente às dimensões informacional e de autoridade do controle humano significativo, de forma na qual, um operador sob efeito desse viés pode dispor formalmente do poder de rejeitar uma recomendação e não o exercer, porque a própria estrutura cognitiva da interação o predispõe a aceitar o *output* algorítmico. Esse ponto

é decisivo para o enquadramento da pesquisa porque manter um humano no circuito não mitiga, por si, o efeito do viés (Johnson, 2021) — a presença humana, longe de operar como salvaguarda automática, pode reproduzir efeitos de redução de esforço e vigilância semelhantes aos observados quando humanos partilham responsabilidades entre si. No caso em estudo, esse mecanismo oferece uma chave interpretativa para o dado de que a verificação humana se limitava, na prática, à confirmação (Faris, 2025) — comportamento compatível com um regime de validação ritualística, e não de exame crítico efetivo. Análises independentes reforçam essa leitura ao documentar depoimentos de operadores segundo os quais o escrutínio individual era, no auge da campanha, reduzido a segundos e tratado como etapa de confirmação, e não de deliberação (Downey, 2025).

O terceiro conceito desloca o foco do operador individual para a estrutura do processo decisório como um todo. Trata-se da cadeia de letalidade — ou *kill chain* — que designa a sequência de etapas que vai da detecção de um alvo até seu engajamento. A IA incide sobre essa cadeia de um modo específico: comprime-a, ao unificar, num mesmo arranjo técnico, a coleta de dados, a identificação do alvo e a recomendação de engajamento, reduzindo drasticamente os intervalos entre observar, avaliar e eliminar. Destaca-se, ainda, a contribuição de Johnson (2021), ao observar que mesmo forças que não obtenham, com a IA, decisões de qualidade superior às humanas auferem vantagem decisiva ao operar com um ciclo decisório encurtado, sobretudo em ambientes que exigem rapidez de decisão através de múltiplos domínios — vantagem que tem como contrapartida o encolhimento da janela disponível para deliberação humana.

Esse “encurtamento do ciclo” é descrito por Rangel e Barbosa (2026), a partir da noção de vigilância letal, como o processo pelo qual observação, decisão e execução passam a operar de forma integrada: ao unificar sensor e arma no mesmo arranjo, a transformação não apenas acelera a decisão sobre o uso da força, mas cria condições para que a ação opere de forma contínua e antecipatória. Embora o trabalho desses autores se desenvolva em registro teórico distinto do adotado aqui, a descrição que oferecem da compressão da cadeia é diretamente útil ao projeto, porque nomeia um mecanismo estrutural — e não apenas o cognitivo — pelo qual o controle humano é pressionado, ou seja, quando a cadeia se comprime, a janela temporal disponível para o julgamento humano encolhe na mesma medida. Destaca-se, ainda, que a literatura sobre o caso reforça esse ponto ao situar os sistemas israelenses no esforço de encurtar o intervalo sensor-a-atirador, considerado, do ponto de vista militar, uma vantagem tática; a aceleração é, portanto, buscada como objetivo, e é essa lógica que explica o interesse dos atores em reduzir a janela disponível para o escrutínio humano (Reichert, 2025).

Com isso, a compressão da *kill chain* articula-se, assim, diretamente com a dimensão temporal do controle humano significativo: o controle exige tempo de deliberação e, se a IA reduz estruturalmente esse tempo, existe uma tensão de fundo entre a lógica de aceleração que justifica o emprego desses sistemas e a exigência de supervisão humana significativa. Essa tensão é, em larga medida, o objeto empírico que o caso Lavender/Gospel permite examinar.

O quarto conceito é o de seletor, e sua introdução exige precisão, porque ele nomeia aquilo que efetivamente circula na cadeia algorítmica de decisão. Um “seletor” é um identificador discreto — um número de telefone, um endereço de e-mail, um conjunto de dados — que funciona como representação operacional de um indivíduo dentro do sistema de vigilância e seleção de alvos (Chamayou, 2015). A partir disso, essa noção é central porque revela a operação de abstração que está no cerne do estabelecimento do alvo: a pessoa não entra no sistema como indivíduo, mas como um feixe de identificadores rastreáveis. O alvo, nesse arranjo, não é o indivíduo, mas o seletor a ele atribuído — atribuição sujeita a determinações errôneas.

A importância analítica do conceito aparece justamente nas evidências sobre a falibilidade dessa atribuição. Rangel e Barbosa (2026) recuperam o relato do ex-analista de inteligência Daniel Hale, que descreveu o grau de incerteza envolvido quando a identificação de suspeitos se baseia em indicadores derivados de metadados. Segundo esse testemunho, há casos em que identificadores são atribuídos erroneamente, descobrindo-se apenas meses ou anos depois que o telefone perseguido como alvo de alto valor pertencia, na verdade, a outra pessoa. Esse tipo de erro expõe, portanto, uma fragilidade estrutural, na qual quando o controle humano se exerce sobre seletores, e não sobre pessoas verificadas, a dimensão informacional do controle significativo fica comprometida na origem — o operador valida o identificador, não um indivíduo. Tal fragilidade é agravada, no caso, pelo fato de as próprias fontes reportarem que o sistema produzia erros de classificação e que seus resultados, ainda assim, eram empregados como base de recomendação (Abraham, 2024; Downey, 2025).

A matriz analítica que sustenta o conceito encontra-se em Chamayou (2015), cuja descrição da construção algorítmica do alvo explicita a operação de abstração em questão. Na “economia do alvo” descrita pelo autor, a vigilância persistente sobrepõe dados sociais, espaciais e temporais para fazer emergir, da massa de informações coletadas sobre um indivíduo, padrões de comportamento — de modo que a atividade registrada passa a operar como alternativa à identidade verificada (Chamayou, 2015). Com isso, o alvo deixa de ser a pessoa e passa a ser o conjunto de identificadores e regularidades a ela atribuídos a partir do cruzamento entre a rede de seus lugares e a de seus vínculos, tornando-se previsível — e

qualquer desvio das regularidades estabelecidas pode disparar o alerta (Chamayou, 2015). É justamente essa conversão do indivíduo em perfil rastreável que o conceito de seletor nomeia — e é sobre ela, e não sobre pessoas diretamente verificadas, que o operador humano efetivamente delibera na cadeia algorítmica.

O quinto conceito — tomada de decisão algorítmica — articula os quatro anteriores em uma noção de regime decisório. Por ele entende-se o regime em que recomendações e predições geradas por sistemas de aprendizado de máquina passam a estruturar — informando, acelerando e ordenando — decisões que permanecem, apenas formalmente, sob autoridade humana. Dessa forma, o traço distintivo desse regime reside menos na substituição do humano pela máquina do que na reconfiguração do papel humano, ou seja, de produtor da decisão, o operador passa a validador de uma decisão pré-estruturada pelo sistema.

É importante, para o rigor do enquadramento, registrar que a própria literatura técnica resiste à imagem da “máquina que mata sozinha”. Análises como a de Johnson (2021) sustentam que o uso predominante da IA militar contemporânea não corresponde a automação plena da decisão de engajamento, mas ao aprimoramento da análise de inteligência e à aceleração da seleção de alvos por meio de melhor fusão de dados e reconhecimento de padrões. A precisão importa porque o problema estratégico que este projeto investiga se situa exatamente na zona intermediária já anunciada na introdução — aquela em que o humano permanece no ciclo, porém sob condições de tempo, informação ou estrutura cognitiva que tornam seu controle progressivamente menos significativo.

A orientação de origem dessa leitura encontra-se em Johnson (2021), que situa a IA militar menos como ruptura que substitui o humano do que como vetor que desloca o plano da decisão estratégica. O autor registra que parte da literatura especula que a IA poderia acelerar o combate ao ponto extremo em que “as ações da máquina superarem o ritmo da tomada de decisão humana e potencialmente deslocarem as bases cognitivas do conflito internacional, desafiando a noção clausewitziana de que a guerra é um empreendimento fundamentalmente humano” (Johnson, 2021, p. 3, tradução nossa); entretanto, frente a esse cenário extremo, o autor pondera que essa aceleração, na verdade, “significa que, no futuro próximo, a tomada de decisão estratégica permanecerá um esforço fundamentalmente humano — ainda que imbuído de graus crescentes de interação com máquinas inteligentes” (Johnson, 2021, p. 173, tradução nossa).

Convergente é a contribuição de Payne (2021), para quem “a IA é uma tecnologia de tomada de decisão, e não uma arma” (Payne, 2021, p. 2, tradução nossa), de forma que, ao longo da existência humana, foram os seres humanos que decidiram as questões pelas quais

lutar e o modo de fazê-lo; agora, decisões sobre o uso da força passam a ser estruturadas também por outros meios. Portanto, é estritamente esse deslocamento que define o escopo aqui examinado: o operador permanece formalmente na decisão, mas as condições cognitivas em que a exerce são reconfiguradas pela presença do sistema. É precisamente a noção de que a IA constitui uma tecnologia de decisão — anunciada na introdução — que sustenta o recorte deste projeto, cujo objeto é a decisão que antecede a arma, e não a arma em si.

6 METODOLOGIA

Esta seção descreve como a pesquisa será construída, as operações concretas que conduzem do problema à avaliação da hipótese, os materiais sobre os quais essas operações incidem e os critérios pelos quais a hipótese será sustentada, qualificada ou refutada. Em termos de desenho, trata-se de pesquisa qualitativa, estruturada como estudo de caso único e baseada em documentação de fontes abertas; adota-se aqui a concepção de estudo de caso como estratégia adequada a perguntas do tipo “como” e “por que” sobre fenômenos contemporâneos sobre os quais o pesquisador tem pouco controle (Yin, 2015), bem como a lógica de rastreamento de processo como forma de examinar o encadeamento causal dentro de um único caso (George; Bennett, 2005). Além dessa caracterização, vale explicitar também o percurso operacional que a materializa — um encadeamento de operações que parte da constituição do corpus documental e chega à avaliação da hipótese, cada qual vinculada aos objetivos específicos que executa.

A opção pelo exame em profundidade de um único caso decorre diretamente do problema, dado que se pergunta de que maneira uma configuração específica — o emprego de dois sistemas concretos, por uma força armada determinada, em período delimitado — alterou o grau de controle humano sobre a decisão de uso da força letal. Dessa forma, o caso viabiliza esse exame por reunir densidade documental em fontes abertas, recorte temporal delimitado e exposição da tensão entre escala algorítmica e supervisão humana. Nesse desenho, a relação investigada é operacionalizada em duas variáveis: a variável dependente é o grau de controle humano sobre a decisão de uso da força letal, aquilo que se busca observar em sua variação; a variável independente é o emprego dos sistemas algorítmicos de apoio à seleção de alvos — Lavender e Gospel — cuja presença, escala e velocidade se supõe condicionarem a primeira. Assim, a investigação consiste, em essência, em examinar como a variação na segunda incide sobre a primeira. Como se trata de caso único, a inferência não repousa sobre comparação entre casos, mas sobre a variação interna ao próprio caso — entre fases da campanha, tipos de alvo e ajustes documentados do limiar de pontuação — e sobre o contraste com a prática de determinação de alvos anterior ao período, tomada como linha de base.

A unidade de análise, por sua vez, é o emprego dos sistemas Lavender e Gospel pelas Forças de Defesa de Israel no conflito contra o Hamas, entre 2023 e 2024. A delimitação é precisa em três eixos: no eixo do objeto, recortam-se especificamente esses dois sistemas, e não a totalidade do aparato de IA militar israelense nem o conflito em sua dimensão geral; no eixo espacial, o contexto é o do conflito em Gaza; no eixo temporal, o recorte cobre 2023–2024, período em que o emprego dos sistemas foi reportado e documentado. Essa delimitação cumpre função operacional, ao definir o que será lido, quais fontes serão buscadas e a fronteira além da qual a evidência deixa de ser pertinente à pesquisa. Convém destacar, ainda, que não se busca a generalização estatística para um universo de conflitos, mas apenas a generalização analítica — a possibilidade de, a partir do caso, refinar e tensionar os conceitos mobilizados, contribuindo para o corpo de conhecimento sobre o controle humano em sistemas de apoio à decisão militar. Portanto, essa ressalva delimita também o alcance da natureza assimétrica do conflito: os achados iluminam o problema do controle humano sobre a geração algorítmica de alvos, sem pretender estender-se automaticamente ao emprego de IA em guerra interestatal convencional, cuja dinâmica de comando, escala e adversário difere e demandaria investigação própria.

Estabelecido o desenho, o percurso operacional inicia-se pela constituição do escopo documental. Nesse sentido, a pesquisa trabalha com documentação secundária, reunida em quatro categorias de fontes abertas: (i) investigações jornalísticas que reconstruíram o funcionamento de Lavender e Gospel a partir de testemunhos de fontes internas e de documentos obtidos pela apuração, das quais provêm os dados centrais que a análise mobiliza — escala de processamento populacional, volume de alvos gerados, sistema de pontuação de probabilidade e tempo médio de verificação humana (Abraham, 2023, 2024); (ii) literatura acadêmica especializada sobre IA militar e seleção algorítmica de alvos, incluindo estudos revisados por pares que analisam diretamente o caso (Downey, 2025; Reichert, 2025), da qual se extraem os conceitos analíticos e as análises já produzidas sobre o caso e sobre casos correlatos; (iii) relatórios de organizações e instituições de pesquisa dedicadas ao estudo da IA militar e dos conflitos contemporâneos, que oferecem contextualização estratégica e dados comparativos; e (iv) documentos públicos e material de inteligência de fontes abertas (OSINT), na medida em que esclareçam o funcionamento e o contexto de emprego dos sistemas. Dois critérios objetivos de inclusão delimitam o escopo: a fonte deve tratar diretamente de Lavender e Gospel, ou dos conceitos aplicáveis à sua análise, e deve situar-se no recorte da unidade de análise. Cada documento incluído será registrado em ficha padronizada — referência completa,

categoria de fonte, data, dados factuais extraídos e conceito a que se vinculam — de modo que o corpus resulte organizado e auditável para as etapas seguintes.

Reunido o corpus, procede-se à qualificação das evidências por verificação cruzada, pois pesquisar um conflito armado em curso impõe assumir que dados provenientes de zonas de conflito são estruturalmente enviesados e incertos — em razão do acesso limitado ao terreno, do uso da informação para fins de propaganda e da indisponibilidade de verificação independente; entidades estatais e não estatais podem difundir deliberadamente informação distorcida, levando à superestimação ou à subestimação tanto da acurácia dos sistemas quanto de seus efeitos.

A fim de mitigar esse risco, adota-se um protocolo objetivo: toda afirmação factual central da análise deverá estar respaldada por, no mínimo, duas fontes independentes, pertencentes a categorias distintas do escopo. Para esse fim, duas fontes só serão consideradas independentes quando não derivarem uma da outra, critério que exige atenção especial, dado que parte da literatura acadêmica sobre o caso reporta os mesmos dados a partir da apuração jornalística original; nesses casos, a corroboração será tratada como parcial, e não plena. Os pronunciamentos oficiais das Forças de Defesa de Israel serão incorporados como fonte de contraditório, de modo que as divergências entre a versão oficial e os relatos de fontes internas fiquem registradas em vez de descartadas. Havendo contradição entre fontes, o caráter controverso da informação será explicitado e nenhum juízo definitivo será firmado sobre base unilateral, priorizando-se a evidência respaldada por elementos verificáveis. Pretende-se, assim, conferir à base de evidências um grau de confiabilidade explícito, que torne acompanhável o trajeto que vai da fonte à conclusão.

Sobre esse escopo já qualificado, reconstrói-se em seguida a cadeia de geração e engajamento de alvos ponto a ponto — coleta e processamento de dados, pontuação, conversão em alvo, rastreamento, autorização do ataque e seleção da munição — descrevendo o funcionamento dos sistemas no grau de precisão que as evidências sustentam. Cada ponto da cadeia é então examinado quanto à participação humana, registrando-se, para cada ponto de intervenção identificado, três informações objetivas: que informação o operador recebia, de quanto tempo dispunha e que poder efetivo de alteração ou veto detinha. Obtém-se, com isso, um mapa da cadeia decisória que localiza, ponto a ponto, onde o controle humano existiu e em que condições foi exercido.

Esse mapa é, na sequência, submetido ao quadro conceitual, traduzindo-se cada uma das três dimensões do controle humano significativo em uma operação concreta sobre as evidências. Na dimensão temporal, confronta-se o tempo de verificação humana documentado

— da ordem de vinte segundos por alvo — e a taxa de geração de alvos no período com o tempo requerido para uma deliberação informada, tomando como parâmetro de comparação o ritmo da construção humana de pastas de alvo na prática anterior das próprias IDF, documentada nas mesmas fontes, e mobilizando o conceito de compressão da cadeia de letalidade para interpretar o resultado. Na dimensão informacional, determina-se sobre o quê o operador exercia controle — um indivíduo verificado ou um seletor derivado de metadados — mapeando os tipos de erro de atribuição relatados nas fontes e avaliando, a partir deles, a integridade da base informacional da decisão. Na dimensão de autoridade efetiva, por fim, classifica-se cada ponto de intervenção humana como escrutínio efetivo ou como ratificação ritualística, segundo critérios observáveis nas fontes — existência ou não de revisão da inteligência bruta, possibilidade real de rejeição da recomendação e frequência com que essa rejeição ocorria — mobilizando o conceito de *automation bias* para interpretar o comportamento documentado dos operadores. Espera-se, assim, uma leitura, dimensão a dimensão, do grau de controle humano efetivamente exercido.

A leitura dimensional é, por fim, confrontada com a hipótese segundo critérios de decisão definidos previamente. Desse modo, a hipótese será considerada corroborada se, em pelo menos duas das três dimensões, evidências sustentadas por fontes independentes indicarem que a intervenção humana se aproximou da ratificação — tempo incompatível com deliberação informada, controle incidindo sobre seletores não verificados e rejeição rara ou sem efeito prático. Por outro lado, será considerada refutada se as evidências indicarem a preservação da deliberação — tempo suficiente, acesso do operador à inteligência subjacente e rejeições frequentes e eficazes. Por último, será qualificada se o quadro se mostrar heterogêneo entre dimensões, fases do conflito ou tipos de alvo, desfecho que exigirá reformular a resposta em termos condicionais, especificando sob quais condições o controle se degrada. Cabe registrar que a dimensão de autoridade efetiva é condição especialmente sensível: caso as evidências indiquem, nela, mera ratificação, o resultado será reportado ainda que as demais dimensões se mostrem preservadas, dado que a capacidade real de veto é o núcleo do controle significativo.

Em qualquer dos três desfechos, o problema de pesquisa terá sido respondido. Cumprida a avaliação, os achados são projetados sobre o quadro mais amplo dos Estudos Estratégicos — compressão do ciclo decisório, doutrina de emprego e distribuição do poder militar — completando a sequência analítica que os quatro objetivos específicos organizam.

7 CRONOGRAMA DE EXECUÇÃO

A execução da pesquisa está prevista para o período correspondente ao curso de mestrado, organizada em etapas que traduzem os objetivos específicos e os procedimentos metodológicos definidos. As fases não são estritamente sequenciais: algumas se sobrepõem,

como o levantamento bibliográfico e a redação, que acompanha progressivamente o avanço analítico. O primeiro semestre concentra-se na consolidação teórica e metodológica; o segundo dedica-se ao levantamento sistemático e à organização crítica das fontes de evidência, com aplicação do protocolo de verificação cruzada; o terceiro constitui o núcleo analítico, com o confronto entre evidências e conceitos estruturado pelos objetivos específicos; e o quarto destina-se à consolidação dos resultados, à redação final e aos procedimentos de defesa. O Quadro 1 apresenta a distribuição das atividades ao longo dos quatro semestres letivos.

Quadro 1 – Cronograma de execução da pesquisa

Atividade	1º sem.	2º sem.	3º sem.	4º sem.
Cumprimento de créditos e disciplinas obrigatórias	X	X		
Aprofundamento do referencial teórico-conceitual	X	X		
Refinamento do desenho de pesquisa e qualificação	X			
Levantamento e organização das fontes de evidência		X	X	
Aplicação do protocolo de verificação cruzada		X	X	
Descrição do funcionamento dos sistemas (obj. a)		X	X	
Identificação da participação humana na cadeia (obj. b)			X	
Análise do efeito de escala e velocidade (obj. c)			X	
Discussão das implicações estratégicas (obj. d)			X	X
Redação dos capítulos			X	X
Revisão final e formatação				X
Depósito e defesa da dissertação				X

Fonte: Elaboração própria (2026).

A previsão admite ajustes decorrentes do calendário do Programa e do andamento do levantamento de fontes. Cabe registrar, em coerência com o caráter de um projeto de pesquisa, que a viabilidade integral do cronograma está condicionada à permanência da disponibilidade de evidências em fontes abertas sobre o caso — condição que, uma vez satisfeita no momento de elaboração do projeto, sustenta o planejamento aqui apresentado.

REFERÊNCIAS

- ABRAHAM, Yuval. **‘A mass assassination factory’: inside Israel’s calculated bombing of Gaza.** +972 Magazine, [s. l.], 30 nov. 2023.
- ABRAHAM, Yuval. **‘Lavender’: the AI machine directing Israel’s bombing spree in Gaza.** +972 Magazine, [s. l.], 3 abr. 2024.
- CHAMAYOU, Grégoire. **Teoria do drone.** São Paulo: Cosac Naify, 2015.
- DORSEY, Jessica. **The erosion of human(e) judgement in targeting? Quantification logics, AI-enabled decision support systems and proportionality assessments in IHL.** International Review of the Red Cross, [s. l.], v. 107, n. 930, p. 1041-1071, 2026. DOI: 10.1017/S1816383125100969.
- DOWNEY, Anthony. **The alibi of AI: algorithmic models of automated killing.** Digital War, [s. l.], v. 6, art. 9, 2025. DOI: 10.1057/s42984-025-00105-7.
- ERSKINE, Toni; MILLER, Steven E. **AI and the decision to go to war: future risks and opportunities.** Australian Journal of International Affairs, [s. l.], v. 78, n. 2, p. 135-147, 2024. DOI: 10.1080/10357718.2024.2349598.
- FARIS, Lamia. **Algorithmic targeting in the Iranian–Israeli confrontation: technical realities, legal thresholds, and the boundaries of human control.** F1000Research, [s. l.], v. 14, n. 1200, 2025. DOI: 10.12688/f1000research.169794.1.
- GEORGE, Alexander L.; BENNETT, Andrew. **Case studies and theory development in the social sciences.** Cambridge, MA: MIT Press, 2005.
- JOHNSON, James. **Artificial intelligence and the future of warfare: the USA, China, and strategic stability.** Manchester: Manchester University Press, 2021.
- PAYNE, Kenneth. **I, warbot: the dawn of artificially intelligent conflict.** London: Hurst & Company, 2021.
- RANGEL, Livia; BARBOSA, Gabriel P. **Comando remoto: drones e a operacionalização da exceção na guerra contemporânea.** Revista Brasileira de Estudos de Defesa, [s. l.], v. 13, e026005, p. 1-23, 2026.
- REICHERT, Ramón. **Autonomous occupation: Israel’s AI-driven drone warfare and the digital architecture of authoritarian power.** Dialogues on Digital Society, [s. l.], v. 1, n. 3, p. 368-372, 2025. DOI: 10.1177/29768640251381423.
- SCHARRE, Paul. **Army of none: autonomous weapons and the future of war.** New York: W. W. Norton & Company, 2018.
- YIN, Robert K. **Estudo de caso: planejamento e métodos.** 5. ed. Porto Alegre: Bookman, 2015.
- ZYSK, Katarzyna. **Struggling, not crumbling: Russian defence AI in a time of war. RUSI Commentary.** London: Royal United Services Institute, 20 nov. 2023.