



Identification and Neutralization of Misinformation in Cognitive Warfare using Machine Learning Algorithms

Reza Taheri^{1✉} / Karim Sadeghi²

1. Corresponding Author, PhD Student in Strategic Knowledge Management, Supreme National Defense University, Tehran, Iran. Email: Taherireza1985@gmail.com

2. Master of Defense Management, Command and Staff College of the Islamic Republic of Iran Army, Tehran, Iran. Email: Karimsadeghi1363@gmail.com

Article Info

Article type:

Research Article

Article history:

Received

06 Jan 2026

Received in revised form

05 Jul 2026

Accepted

21 Jul 2026

Published online

31 Jul 2026

Keywords:

Cognitive Warfare, Information Disorders, Machine Learning, Cognitive Resilience, Decision Support Systems.

ABSTRACT

Cognitive warfare aims to weaken the ability to make correct decisions by targeting the minds and perceptions of commanders. The spread of false and misleading information is among the most critical tools of this warfare used to create confusion and disruption in the military decision-making process. This research aims to answer how machine learning algorithms can be utilized to identify and neutralize misinformation in cognitive warfare. The research methodology is qualitative, based on a systematic literature review and thematic analysis, through which reliable national and international sources have been analyzed.

The findings indicate that machine learning algorithms, across three layers—content analysis, network analysis, and user behavior analysis—are capable of identifying various types of threats, including fake news, propaganda, rumors, and conspiracy theories. Based on the results, deep learning, natural language processing (NLP), and traditional machine learning algorithms can neutralize the destructive effects of cognitive warfare through operational solutions such as automatic removal, content labeling, and providing correct information. Finally, this research emphasizes that despite the capabilities of artificial intelligence, maintaining “human-in-the-loop” interaction is essential to prevent cognitive laziness and ensure compliance with ethical and legal considerations on the battlefield.

Cite this article: Taheri, Reza. Sadeghi, Karim. (2026). Identification and Neutralization of Misinformation in Cognitive Warfare using Machine Learning Algorithms. *Scientific Journal of Cognitive Warfare Sciences and Technologies*, 56 (1), 1-20.

Publisher: AJA University of Command and Staff,

© “Authors retain the copyright and full publishing rights.”

DOI: <https://doi.org/10.22034/jcwst.2026.571146.1026>





شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ شناختی با استفاده از

الگوریتم‌های یادگیری ماشین

رضا طاهری^۱ | کریم صادقی^۲

۱. نویسنده مسئول، دانشجوی مقطع دکتری مدیریت راهبردی دانش دانشگاه عالی دفاع ملی، تهران، ایران رایانامه:

Taherireza1985@gmail.com

۲. کارشناس ارشد مدیریت دفاعی دافوس آجا، تهران، ایران رایانامه: Karimsadeghi1363@gmail.com

اطلاعات مقاله	چکیده
نوع مقاله:	جنگ‌شناختی با هدف قرار دادن ذهن و ادراک فرماندهان، توانایی تصمیم‌گیری
مقاله پژوهشی	صحیح را تضعیف می‌کند. انتشار اطلاعات نادرست و گمراه‌کننده از مهم‌ترین ابزارهای
تاریخچه مقاله:	این نبرد برای ایجاد سردرگمی و اختلال در فرآیند تصمیم‌گیری نظامی است. این
تاریخ دریافت:	پژوهش با هدف پاسخ به این سؤال انجام شده است که چگونه می‌توان با استفاده از
تاریخ بازنگری:	الگوریتم‌های یادگیری ماشین، اطلاعات نادرست را در جنگ‌شناختی شناسایی و
تاریخ پذیرش:	خنثی کرد. روش تحقیق حاضر، کیفی و بر مبنای مرور نظام‌مند ادبیات و تحلیل
تاریخ انتشار:	مضمون است که طی آن منابع معتبر داخلی و بین‌المللی واکاوی شده‌اند. یافته‌های
کلیدواژه‌ها:	پژوهش نشان داد که الگوریتم‌های یادگیری ماشین در سه لایه تحلیل محتوا، تحلیل
جنگ‌شناختی، تاب‌آوری	شبکه و تحلیل رفتار کاربر قادر به شناسایی گونه‌های مختلف تهدید شامل اخبار
شناختی، سامانه‌های	جعلی، پروپاگاندا، شایعات و نظریه‌های توطئه هستند. بر اساس نتایج، الگوریتم‌های
پشتیبانی تصمیم،	یادگیری عمیق، پردازش زبان طبیعی و یادگیری ماشین سنتی می‌توانند از طریق
اختلالات اطلاعاتی،	راهکارهای عملیاتی نظیر حذف خودکار، برچسب‌گذاری محتوا و ارائه اطلاعات
یادگیری ماشین	صحیح، اثرات مخرب جنگ‌شناختی را خنثی کنند. درنهایت، این پژوهش تأکید
	می‌کند که علی‌رغم توانمندی‌های هوش مصنوعی، حفظ تعامل «انسان در حلقه»
	برای جلوگیری از تنبلی شناختی و رعایت ملاحظات اخلاقی و حقوقی در صحنه نبرد
	ضروری است.

استناد: طاهری، رضا؛ صادقی، کریم (۱۴۰۴). شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ‌شناختی با استفاده از الگوریتم‌های

یادگیری ماشین. فصلنامه علمی علوم و فناوری‌های جنگ شناختی، ۲ (۴)، ۲۰-۱.

DOI: <http://doi.org/10.22034/jcwst.2026.571146.1026>

ناشر: دانشگاه فرماندهی و ستاد ارتش جمهوری اسلامی ایران

مقدمه:

در فضای نبردهای آینده، پیچیدگی و سرعت بالا، تصمیم‌گیری را با تکیه صرف بر توانمندی‌های انسانی دشوار می‌سازد (افشردی و شیخ، ۱۳۹۸). به گفته رضایی، رشید و پوردستان (۱۳۹۹)، فرماندهی و کنترل هوشمند به‌عنوان یک رویکرد نوین در صحنه نبرد، می‌تواند به فرماندهان در مقابله با این پیچیدگی و سرعت کمک کند (رضایی و همکاران، ۱۳۹۹). یکی از مهم‌ترین وظایف فرماندهی و کنترل هوشمند، تشخیص و مقابله با تهدیدات اطلاعاتی است که اطلاعات نادرست و گمراه‌کننده یکی از مهم‌ترین آن‌ها به شمار می‌رود (شاملو، ۱۴۰۱). در دنیای امروز، با گسترش فناوری‌های اطلاعات و ارتباطات و ظهور رسانه‌های اجتماعی، جنگ به شکلی پیچیده‌تر و نامحسوس‌تر از گذشته در آمده است (لوگویادر^۱، ۲۰۲۲). جنگ‌شناختی به‌عنوان یکی از ابعاد این جنگ جدید، با هدف تأثیرگذاری بر افکار، باورها و رفتارهای افراد و گروه‌ها، به‌ویژه در حوزه نظامی، مطرح شده است (کلاوری و دو کلوزل^۲، ۲۰۲۲). یکی از مهم‌ترین ابزارهای مورد استفاده در جنگ‌شناختی، انتشار اطلاعات نادرست و گمراه‌کننده است که می‌تواند باعث ایجاد سردرگمی، کاهش اعتماد به نفس و اختلال در فرآیند تصمیم‌گیری در فرماندهان و نیروهای نظامی شود (میلر^۳، ۲۰۲۳). با توجه به اهمیت تشخیص و خنثی‌سازی اطلاعات نادرست در جنگ‌شناختی، روش‌های مختلفی برای این کار ارائه شده است. یکی از روش‌های نوین و مؤثر در این زمینه، استفاده از الگوریتم‌های یادگیری ماشین است (بووارپ^۴، ۲۰۲۳). الگوریتم‌های یادگیری ماشین می‌توانند با تحلیل داده‌های بزرگ و شناسایی الگوهای مخفی در آن‌ها، به‌طور خودکار اطلاعات نادرست را تشخیص دهند و به فرماندهان و نیروهای نظامی در مقابله با آن‌ها کمک کنند (شاملو، ۱۴۰۱). در این پژوهش، به بررسی این سؤال می‌پردازیم که «چگونه می‌توان با استفاده از الگوریتم‌های یادگیری ماشین، اطلاعات نادرست و گمراه‌کننده را در جنگ‌شناختی شناسایی و خنثی کرد؟» برای پاسخ به این سؤال، ابتدا به بررسی مفاهیم کلیدی مانند جنگ‌شناختی، اطلاعات نادرست و الگوریتم‌های یادگیری ماشین می‌پردازیم. سپس، مطالعات انجام شده

¹ Le Guyader² Claverie & Du Cluzel³ Miller⁴ Buvarp

در این زمینه را مرور می‌کنیم و درنهایت، یک چارچوب برای شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ‌شناختی با استفاده از الگوریتم‌های یادگیری ماشین ارائه می‌دهیم.

مرور پیشینه و مبانی نظری

پیشینه و سابقه پژوهش

در این بخش، به بررسی پیشینه پژوهش‌های داخلی و خارجی در حوزه شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ‌شناختی با استفاده از الگوریتم‌های یادگیری ماشین می‌پردازیم؛ لازم به ذکر است این مطالعات بر اساس اهمیت و ارتباط آن‌ها با موضوع پژوهش مرتب گردیده‌اند.

مطالعات داخلی

در پژوهشی که توسط فرهنگ و همکاران در سال ۱۴۰۲ با عنوان «شناسایی اخبار جعلی در شبکه‌های اجتماعی با استفاده از یادگیری عمیق» انجام شد، محققین به بررسی استفاده از یادگیری عمیق برای شناسایی اخبار جعلی در شبکه‌های اجتماعی پرداختند. نتایج پژوهش نشان داد که یادگیری عمیق می‌تواند با دقت بالایی اخبار جعلی را در شبکه‌های اجتماعی تشخیص دهد (فرهنگ و آروند، ۱۴۰۲). ترسلی طی پژوهشی با عنوان «کاربرد هوش مصنوعی در تشخیص اخبار جعلی» در سال ۱۴۰۳ به بررسی چگونگی استفاده از مدل‌های یادگیری ماشین در پیش‌بینی و تشخیص اخبار جعلی پرداخت. نتایج پژوهش نشان داد که مدل یادگیری ماشین جنگل تصادفی و مدل یادگیری تقویت گرادیان بهترین نتیجه را در تشخیص اخبار جعلی دارند. این یافته‌ها می‌توانند در انتخاب الگوریتم‌های مناسب برای شناسایی اطلاعات نادرست در جنگ‌شناختی مورد استفاده قرار گیرند (ترسلی، ۱۴۰۳).

مداح و شاه‌محمدی طی پژوهشی در سال ۱۴۰۱ با عنوان «فرصت‌ها و چالش‌های هوش مصنوعی در جنگ‌شناختی» به بررسی نقش هوش مصنوعی در جنگ‌شناختی پرداختند. آن‌ها در این پژوهش به بررسی فرصت‌ها و چالش‌های هوش مصنوعی در جنگ‌شناختی پرداخته و به این نتیجه رسیدند که هوش مصنوعی می‌تواند در تشخیص تهدیدات، پیش‌بینی آینده، پاسخ هوشمند و مدیریت داده‌ها در جنگ‌شناختی مؤثر باشد (مداح و شاه‌محمدی، ۱۴۰۲). در مقاله «هوش مصنوعی و تأثیر آن بر امنیت سایبری و حفاظت از داده‌ها» که توسط محمودی و بحرکاظمی در سال ۱۴۰۱ منتشر شد، محققین به

بررسی نقش هوش مصنوعی در امنیت سایبری و حفاظت از داده‌ها پرداختند. یافته‌های این پژوهش نشان می‌دهد که هوش مصنوعی می‌تواند به سازمان‌ها در حفظ امنیت سایبری خود و تشخیص و پیشگیری از تهدیدات سایبری کمک کند. (محمودی و بحرکاظمی، ۱۴۰۳)

در پژوهشی که توسط شاملو در سال ۱۴۰۱ با عنوان «جایگاه فناوری هوش مصنوعی در سامانه‌های فرماندهی و کنترل آینده (مطالعه موردی: سامانه فرماندهی و کنترل فراگیر مشترک آمریکا)» انجام شد، محقق به بررسی جایگاه فناوری هوش مصنوعی در سامانه‌های فرماندهی و کنترل آینده پرداخته است. یافته‌های این پژوهش نشان می‌دهد که هوش مصنوعی می‌تواند نقش مهمی در ارتقاء سرعت، دقت و اثربخشی تصمیم‌گیری در سامانه‌های فرماندهی و کنترل آینده داشته باشد (شاملو، ۱۴۰۱). در مطالعه‌ای که توسط افشردی و شیخ در سال ۱۳۹۸ با عنوان «بررسی نیازها، قابلیت‌ها و توانمندی‌های فرماندهی و مدیریت در عرصه جنگ‌های آینده ج.ا.ایران» انجام شد، محققین به بررسی نیازمندی‌ها و توانمندی‌های فرماندهی و مدیریت در جنگ‌های آینده پرداختند. یافته‌های این پژوهش نشان می‌دهد که هوش مصنوعی می‌تواند در تقویت توانمندی‌های شناختی فرماندهان و مقابله با اطلاعات نادرست در جنگ‌شناختی مؤثر باشد. این یافته‌ها می‌توانند در طراحی و پیاده‌سازی سیستم‌های مبتنی بر هوش مصنوعی برای مقابله با اطلاعات نادرست در جنگ‌شناختی مورد استفاده قرار گیرند (افشردی و شیخ، ۱۳۹۸).

رضایی، رشید و پوردستان در سال ۱۳۹۹ طی پژوهشی با عنوان «مؤلفه‌ها و ویژگی‌های فرماندهی و کنترل هوشمند در صحنه نبرد» به بررسی اهمیت فرماندهی و کنترل هوشمند در صحنه نبرد پرداختند. با استفاده از روش توصیفی-تحلیلی به بررسی مؤلفه‌ها و ویژگی‌های فرماندهی و کنترل هوشمند در صحنه نبرد پرداخته‌اند. یافته‌های این پژوهش نشان می‌دهد که فرماندهی و کنترل هوشمند می‌تواند نقش مهمی در ارتقاء عملکرد نیروهای نظامی در صحنه نبرد داشته باشد. (رضایی و همکاران، ۱۳۹۹). بختیاری و عسکری در سال ۱۴۰۱ در مطالعه‌ای با عنوان «ارائه مدل آگاهی وضعیتی در شبکه فرماندهی و کنترل پدافند هوایی» به بررسی اهمیت آگاهی وضعیتی در شبکه فرماندهی و کنترل پدافند هوایی پرداختند. با استفاده از روش توصیفی-پیمایشی به بررسی اهمیت آگاهی وضعیتی در شبکه فرماندهی و کنترل پدافند هوایی پرداخته‌اند. یافته‌های این پژوهش نشان می‌دهد که آگاهی وضعیتی می‌تواند نقش مهمی در ارتقاء عملکرد

فرماندهان در تصمیم‌گیری و کنترل عملیات داشته باشد. با توجه به اینکه سیستم‌های پشتیبانی تصمیم مبتنی بر هوش مصنوعی می‌توانند به فرماندهان در افزایش آگاهی موقعیتی و درک بهتر از محیط عملیاتی کمک کنند، می‌توان نتیجه گرفت که این سیستم‌ها می‌توانند به‌طور مؤثری در ارتقاء تاب‌آوری شناختی فرماندهان در جنگ شناختی نقش ایفا کنند (بختیاری و عسکری، ۱۴۰۱). کنعانی و یآوری (۱۴۰۱) در مطالعه‌ای با عنوان «طراحی مدلی برای فرماندهی و کنترل در پلیس هوشمند» به بررسی نقش هوش مصنوعی در ارتقاء عملکرد فرماندهی و کنترل در پلیس هوشمند پرداختند. یافته‌های این پژوهش نشان می‌دهد که هوش مصنوعی می‌تواند با ارائه ابزارهای هوشمند برای جمع‌آوری، تحلیل و تفسیر اطلاعات، شناسایی الگوها و روندها و اتخاذ تصمیمات بهینه در شرایط پیچیده و پویا، به فرماندهان در پلیس هوشمند کمک کند. (کنعانی و یآوری، ۱۴۰۱)

مطالعات خارجی

مطالعات خارجی عبدالعزیز و مروان البحر (۲۰۲۰) در پژوهشی با عنوان «یک مقایسه تجربی از تشخیص اخبار جعلی با استفاده از الگوریتم‌های مختلف یادگیری ماشین» به بررسی عملکرد الگوریتم‌های مختلف یادگیری ماشین در تشخیص اخبار جعلی پرداختند. آن‌ها به این نتیجه رسیدند که الگوریتم بیز ساده در تشخیص اخبار جعلی به‌طور قابل‌توجهی بهتر از سایر الگوریتم‌ها عمل می‌کند (البحر و البحر، ۲۰۲۰). هو و همکاران (۲۰۲۴) در پژوهشی با عنوان «مروری بر تشخیص اخبار جعلی از دیدگاهی جدید» به بررسی روش‌های مختلف تشخیص اخبار جعلی و ارائه چارچوبی جدید برای این کار پرداختند. آن‌ها به این نتیجه رسیدند که روش‌های مختلف تشخیص اخبار جعلی را می‌توان بر اساس ویژگی‌های آن‌ها به سه دسته تقسیم کرد: روش‌های مبتنی بر ویژگی‌های عمدی، روش‌های مبتنی بر انتشار و روش‌های مبتنی بر موضع‌گیری. (هو و همکاران، ۲۰۲۴). در سال ۲۰۱۸، کاترینا کرتیسووا^۱ در مقاله‌ای با عنوان «هوش مصنوعی و اطلاعات نادرست: چگونه هوش مصنوعی روش تولید، انتشار و مقابله با اطلاعات نادرست را تغییر می‌دهد» به بررسی چالش‌ها و فرصت‌های هوش مصنوعی در زمینه اطلاعات نادرست پرداخت. او به این نتیجه رسید که هوش مصنوعی می‌تواند به‌طور مؤثری برای

¹ Kertysova

شناسایی و مقابله با اطلاعات نادرست مورد استفاده قرار گیرد، اما درعین حال، می‌تواند برای تولید و انتشار اطلاعات نادرست نیز مفید باشد. (کرتیسووا، ۲۰۱۸). در سال ۲۰۲۱، کاترینا شدووا^۱ و همکاران در مرکز امنیت و فناوری‌های نوظهور^۲ طی پژوهشی با عنوان «هوش مصنوعی و آینده کمپین‌های اطلاعات نادرست: بخش دوم، مدل تهدید» به بررسی چگونگی تأثیر هوش مصنوعی و یادگیری ماشین بر کمپین‌های اطلاعات نادرست پرداختند. آن‌ها به این نتیجه رسیدند که هوش مصنوعی می‌تواند به‌طور مؤثری برای تولید و انتشار اطلاعات نادرست در مقیاس وسیع مورد استفاده قرار گیرد، اما درعین حال، می‌تواند برای شناسایی و مقابله با اطلاعات نادرست نیز مفید باشد (شدووا، ۲۰۲۱). در سال ۲۰۲۰، آلیم آل ایوب احمد^۳ و همکاران طی پژوهشی با عنوان «تشخیص اخبار جعلی با استفاده از یادگیری ماشین: مروری نظام‌مند بر منابع» به بررسی کاربردهای یادگیری ماشین در تشخیص اخبار جعلی پرداختند. آن‌ها به بررسی انواع طبقه‌بندی کننده‌های یادگیری ماشین که برای تشخیص اخبار جعلی استفاده می‌شوند، روش‌های آموزش و چالش‌های استفاده از یادگیری ماشین در تشخیص اخبار جعلی پرداختند (احمد و همکاران، ۲۰۲۱). سوئانتیرادوی پی^۴ و همکاران (۲۰۲۰) در پژوهشی با عنوان «تشخیص پروپاگاندا از مقالات خبری با استفاده از یادگیری عمیق» به بررسی استفاده از یادگیری عمیق، به‌ویژه شبکه‌های عصبی کانولوشن، برای تشخیص پروپاگاندا در مقالات خبری پرداختند. آن‌ها به این نتیجه رسیدند که شبکه‌های عصبی کانولوشن می‌توانند با دقت بالایی پروپاگاندا را در مقالات خبری تشخیص دهند (سوئانتیرادوی پی و همکاران، ۲۰۲۰).

در سال ۲۰۲۳، آندرا ساندو^۵ و همکاران طی پژوهشی با عنوان «کاربردهای یادگیری ماشین و یادگیری عمیق در تشخیص اطلاعات نادرست: ارزیابی کتاب‌سنجی» به بررسی مقالاتی که از الگوریتم‌های یادگیری ماشین و یادگیری عمیق برای تشخیص اطلاعات نادرست استفاده می‌کنند، پرداختند. آن‌ها به این نتیجه رسیدند که علاقه به موضوع اطلاعات نادرست در جامعه علمی در حال افزایش است و روش‌های یادگیری ماشین و

¹ Sedova

² CSET

³ Ahmed et al

⁴ Sothanthirathlevi P et a

⁵ Sandu et al

یادگیری عمیق می‌توانند به‌طور مؤثری در تشخیص اطلاعات نادرست مورد استفاده قرار گیرند (ساندو و همکاران، ۲۰۲۴). در سال ۲۰۲۴، کاستوپولوس^۱ داوازوس و کوتسیانیتیس به بررسی سیستم‌های پشتیبانی تصمیم مبتنی بر هوش مصنوعی قابل توضیح پرداختند. آن‌ها در پژوهش خود با عنوان «سیستم‌های پشتیبانی تصمیم مبتنی بر هوش مصنوعی قابل توضیح: مروری جدید» به این نتیجه رسیدند که این سیستم‌ها می‌توانند با ارائه توضیحات شفاف و قابل فهم در مورد تصمیمات خود، به افزایش اعتماد کاربران و بهبود فرآیند تصمیم‌گیری کمک کنند. این یافته‌ها می‌تواند به افزایش اعتماد فرماندهان در مواجهه با اطلاعات نادرست کمک کند (کاستوپولوس و همکاران، ۲۰۲۴). شوبرت^۲ و همکارانش در سال ۲۰۱۸ طی پژوهشی با عنوان «هوش مصنوعی برای پشتیبانی تصمیم در سیستم‌های فرماندهی و کنترل» به بررسی نقش هوش مصنوعی در سیستم‌های پشتیبانی تصمیم در حوزه نظامی پرداختند. آن‌ها به این نتیجه رسیدند که هوش مصنوعی می‌تواند با ارائه اطلاعات دقیق و به‌موقع، شناسایی و خنثی‌سازی اطلاعات نادرست و شبیه‌سازی سناریوهای مختلف، به فرماندهان در درک بهتر موقعیت کمک کند (شوبرت و همکاران، ۲۰۱۸).

در سال ۲۰۲۳، بووارپ در پژوهشی با عنوان «گزارش ارزیابی فنی سمپوزیوم تحقیقاتی آر اس وای^۳ کاهش و پاسخ به جنگ‌شناختی» به بررسی موضوع جنگ‌شناختی و روش‌های مقابله با آن پرداخت. یافته‌های این پژوهش نشان می‌دهد که استفاده از هوش مصنوعی و الگوریتم‌های یادگیری ماشین می‌تواند به‌طور مؤثری در شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ‌شناختی مورد استفاده قرار گیرد (بووارپ، ۲۰۲۳). در سال ۲۰۱۸، ون دن بوش و برونکهورست^۴ طی پژوهشی با عنوان «همکاری انسان و هوش مصنوعی برای بهبود تصمیم‌گیری نظامی» به بررسی نقش هوش مصنوعی در تصمیم‌گیری نظامی پرداختند. آن‌ها به این نتیجه رسیدند که همکاری انسان و هوش مصنوعی می‌تواند به‌طور مؤثری در بهبود فرآیند تصمیم‌گیری نظامی و افزایش تاب‌آوری شناختی فرماندهان در جنگ‌شناختی مورد استفاده قرار گیرد (ون دن بوش و

^۱ Kostopoulos et al

^۲ Schubert et al

^۳ RSY

^۴ Van Den Bosch & Bronkhorst

برونکهورست، ۲۰۱۸). در سال ۲۰۲۴، نادیبایدزه^۱، بود و ژانگ در پژوهشی با عنوان «هوش مصنوعی در سیستم‌های پشتیبانی تصمیم نظامی: مروری بر تحولات و مباحث» به این نتیجه رسیدند که هوش مصنوعی می‌تواند با شناسایی و خنثی‌سازی اطلاعات نادرست به فرماندهان در اتخاذ تصمیمات بهینه کمک کند (نادیبایدزه و همکاران، ۲۰۲۴). گریپل^۲، دیویسون و هیندز (۲۰۲۴) در پژوهشی با عنوان «هوش مصنوعی و فناوری‌های مرتبط در تصمیم‌گیری نظامی در مورد استفاده از قدرت در درگیری‌های مسلحانه» به بررسی چالش‌ها و خطرات مرتبط با سیستم‌های پشتیبانی تصمیم مبتنی بر هوش مصنوعی پرداخته‌اند. آن‌ها به این نتیجه رسیدند که شفافیت، تفسیرپذیری و قابلیت پیش‌بینی ازجمله چالش‌های اصلی هستند (گریپل و همکاران، ۲۰۲۴). در سال ۲۰۲۴، پاولیکی و همکاران طی پژوهشی با عنوان «بینش‌های پیشرفته از طریق تحلیل نظام‌مند: ترسیم مسیرهای تحقیقاتی آینده و فرصت‌ها برای XAI^۳ در یادگیری عمیق و هوش مصنوعی مورد استفاده در امنیت سایبری» به این نتیجه رسیدند که هوش مصنوعی قابل توضیح می‌تواند نقش مهمی در ارتقاء تاب‌آوری شناختی فرماندهان داشته باشد (پاولیکی و همکاران، ۲۰۲۴). در سال ۲۰۲۰، مورگان^۴ و همکاران در مطالعه‌ای با عنوان «کاربردهای نظامی هوش مصنوعی» به این نتیجه رسیدند که هوش مصنوعی می‌تواند با تجزیه و تحلیل داده‌های حجیم به فرماندهان در تصمیم‌گیری بهتر و سریع‌تر کمک کرده و تاب‌آوری شناختی آن‌ها را افزایش دهد (مورگان و همکاران، ۲۰۲۰). در سال ۲۰۲۲، لو گویدار طی پژوهشی با عنوان «حوزه شناختی: ششمین حوزه عملیات» به این نتیجه رسید که حوزه شناختی به‌طور فزاینده‌ای در جنگ‌های مدرن اهمیت پیدا کرده است (لو گویدار، ۲۰۲۲). زینگدونگ و زیانگمینگ طی پژوهشی در سال ۲۰۲۲ با عنوان «جنگ‌شناختی الگوریتمی: تغییر پارادایم جنگ افکار عمومی در زمینه مناقشه روسیه و اوکراین» به این نتیجه رسیدند که الگوریتم‌ها می‌توانند به‌طور مؤثری برای انتشار اطلاعات نادرست مورد استفاده قرار گیرند (زینگدونگ و زیانگمینگ، ۲۰۲۲). در نهایت، علویطبی و الدوساری (۲۰۲۴) در پژوهشی با عنوان «مروری بر تکنیک‌های تشخیص اخبار جعلی

¹ Nadibaidze et al

² Greipl et al

³ Explainable Artificial Intelligence

⁴ Morgan et al

برای زبان عربی» به بررسی روش‌های مختلف تشخیص اخبار جعلی و چالش‌های موجود در زبان عربی پرداختند (علویطبی و الدوساری، ۲۰۲۴).

مرور ادبیات

در چشم‌انداز نبردهای آینده، پیچیدگی و سرعت بالای تحولات، تصمیم‌گیری را با تکیه صرف بر توانمندی‌های انسانی دشوار می‌سازد (افشردی و شیخ، ۱۳۹۸). با ظهور فناوری‌های نوین اطلاعاتی و ارتباطی، ماهیت جنگ از تقابل‌های سخت‌افزاری فراتر رفته و عرصه‌های جدیدی را در بر گرفته است (لوگویادر، ۲۰۲۲). در این زیست‌بوم جدید، تسلط بر جریان اطلاعات و حفاظت از ذهن فرماندهان در برابر فریب، به مؤلفه‌ای تعیین‌کننده تبدیل شده است. این بخش به تبیین مفاهیم بنیادین پژوهش شامل ماهیت جنگ شناختی، گونه‌شناسی اطلاعات نادرست و نقش الگوریتم‌های هوشمند ارتقاء تاب‌آوری سامانه فرماندهی و کنترل می‌پردازد.

۱. جنگ شناختی: جنگ شناختی که اکنون به‌عنوان «ششمین حوزه عملیات» شناخته می‌شود، نشان‌دهنده تغییر ماهیت جنگ‌ها به سمت عرصه‌های غیرفیزیکی است (لوگویادر، ۲۰۲۲). برخلاف نبردهای کلاسیک که هدفشان تصرف سرزمین بود، هدف غایی در جنگ شناختی، نفوذ به لایه‌های عمیق ذهن و تغییر درک، باورها و رفتار فرماندهان و نیروهای نظامی است (کلاوری و دوکلوزل، ۲۰۲۲). دشمن در این بستر با بهره‌گیری از ابزارهایی نظیر عملیات روانی، تبلیغات و نفوذ سایبری، سعی در تضعیف روحیه، ایجاد تفرقه و اختلال در فرآیندهای حیاتی تصمیم‌گیری دارد (فنسترمکر و همکاران، ۲۰۲۳). در چنین شرایطی، ارتقاء تاب‌آوری شناختی فرماندهان برای حفظ ثبات در تصمیم‌گیری، ضرورتی انکارناپذیر است (رضایی و همکاران، ۱۳۹۹).

۲. گونه‌شناسی اختلالات اطلاعاتی در صحنه نبرد: سلاح اصلی در جنگ شناختی، انتشار هدفمند اطلاعات نادرست و گمراه‌کننده است. این اطلاعات که می‌توانند به‌صورت متن، تصویر یا ویدئو باشند، با هدف ایجاد سردرگمی و کاهش اعتمادبه‌نفس در شبکه فرماندهی منتشر می‌شوند (هو و همکاران، ۲۰۲۴). بر اساس ادبیات پژوهش، مهم‌ترین مصادیق این تهدیدات عبارت‌اند از:

- اخبار جعلی: اطلاعاتی کاملاً ساختگی و کذب که با ظاهری مشابه اخبار معتبر طراحی می‌شوند تا مخاطب را فریب دهند. مطالعات نشان می‌دهد که

الگوریتم‌های یادگیری ماشین نظیر «بیز ساده» در تشخیص این اخبار عملکرد قابل توجهی دارند (البحر و البحر، ۲۰۲۰).

- **پروپاگاندا:** محتوایی که لایه‌هایی از حقیقت را با مبالغه و گزینش‌گری ترکیب می‌کند (سوئانتیرادوی و همکاران، ۲۰۲۰). ترکیب صدق و کذب در پروپاگاندا باعث می‌شود قدرت نفوذ آن در ذهن فرماندهان بیشتر باشد و تشخیص آن نیازمند تحلیل‌های معنایی پیچیده است.

- **شایعات و نظریه‌های توطئه:** ادعاهای تأیید نشده یا نظام‌مندی که با هدف قرار دادن سوگیری‌های شناختی، اعتماد به ساختار فرماندهی را هدف می‌گیرند و منجر به فلج تصمیم‌گیری در لحظات حساس می‌شوند (هو و همکاران، ۲۰۲۴).

۳. **سیستم‌های پشتیبانی تصمیم:** سیستم‌های پشتیبانی تصمیم مبتنی بر هوش مصنوعی، با هدف تقویت توانمندی‌های انسانی و نه جایگزینی آن‌ها طراحی شده‌اند (شوبرت و همکاران، ۲۰۱۸). این سیستم‌ها با ارائه اطلاعات دقیق، شبیه‌سازی سناریوها و پیش‌بینی تهدیدات، آگاهی موقعیتی فرماندهان را ارتقاء می‌دهند (شاملو، ۱۴۰۱). با این حال، ادبیات پژوهش تأکید دارد که شناسایی و خنثی‌سازی نباید به‌طور کاملاً خودمختار به ماشین واگذار شود. رویکرد انسان در حلقه بیان می‌کند که هوش مصنوعی باید نقش دستیار را ایفا کند و قضاوت نهایی، به‌ویژه در تصمیمات حساس نظامی که پیامدهای اخلاقی دارد، باید بر عهده انسان باشد (ون دن بوش و برونکهورست، ۲۰۱۸). اتکای صرف به ماشین می‌تواند منجر به تنبلی شناختی و کاهش قدرت تحلیل مستقل فرمانده شود. همچنین، استفاده از هوش مصنوعی قابل توضیح با شفاف‌سازی علل تصمیم‌گیری الگوریتم، اعتماد فرماندهان را به سیستم افزایش داده و ابهام را در شرایط بحرانی کاهش می‌دهد (کاستوپولوس و همکاران، ۲۰۲۴).

متغیرهای پژوهش

- **متغیر مستقل:** الگوریتم‌های یادگیری ماشین (انواع، ویژگی‌ها، عملکرد)
- **متغیر وابسته:** شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ شناختی (روش‌ها، معیارها، اثربخشی)
- **متغیرهای تعدیل‌کننده:** نوع اطلاعات نادرست (متن، تصویر، ویدئو)، بستر انتشار (شبکه‌های اجتماعی، وبسایت‌ها)، اهداف جنگ شناختی (تضعیف روحیه، ایجاد تردید، اختلال در تصمیم‌گیری)

روش‌شناسی

با توجه به ماهیت اکتشافی و تحلیلی این پژوهش که در پی واکاوی عمیق «چگونه می‌توان با استفاده از الگوریتم‌های یادگیری ماشین، اطلاعات نادرست و گمراه‌کننده را در جنگ‌شناختی شناسایی و خنثی کرد؟ است، رویکرد کیفی برای انجام این پژوهش انتخاب شد. در این راستا، از روش توصیفی-تحلیلی و مرور نظام‌مند ادبیات به‌عنوان ابزار اصلی گردآوری و تحلیل داده‌ها استفاده شده است. مرور نظام‌مند ادبیات به پژوهشگران امکان می‌دهد تا به شکلی جامع و ساختارمند، دانش موجود در حوزه مورد مطالعه را بررسی و ارزیابی کنند و به کشف الگوها، روندها و روابط بین مفاهیم کلیدی بپردازند.

پروتکل پژوهش

الف- شناسایی و انتخاب منابع: برای انجام مرور نظام‌مند، جستجوی گسترده‌ای در پایگاه‌های اطلاعاتی داخلی مانند پایگاه اطلاعات علمی جهاد دانشگاهی (اس‌آی‌دی)^۱، مرکز اطلاعات علمی ایران (ایرانداک) و پایگاه مجلات تخصصی نور (نورمگز) و پایگاه‌های اطلاعاتی بین‌المللی مانند اسکوپوس^۲، وب‌آو‌ساینس^۳ و گوگل اسکولار^۴ انجام شد. در این مرحله، از کلیدواژه‌های زیر برای جستجو در پایگاه‌های اطلاعاتی استفاده شد: کلیدواژه‌های فارسی: جنگ‌شناختی، اطلاعات نادرست، الگوریتم‌های یادگیری ماشین، هوش مصنوعی، شناسایی، "نشی‌ازی".

کلیدواژه‌های انگلیسی

"("Cognitive Warfare" OR "Information Warfare") AND ("Disinformation" OR "Misinformation" OR "Fake News") AND ("Machine Learning Algorithms" OR "Artificial Intelligence" OR "AI") AND ("Detection" OR "Neutralization").

ب- غربالگری و انتخاب نهایی: غربالگری اولیه پس از ثبت مقالات در نرم‌افزار استناددهی اندنوت^۵، بر اساس عنوان و چکیده انجام شد. در این مرحله، مقالاتی که به‌طور مستقیم به موضوع پژوهش مرتبط نبودند و یا کیفیت پایینی داشتند، حذف شدند. در ادامه، تمام متن مقالات باقی‌مانده به‌دقت بررسی شد. معیارهای ارزیابی شامل ارتباط با سؤال

^۱ SID

^۲ Scopus

^۳ Web of Science

^۴ Google Scholar

^۵ EndNote

پژوهش، کیفیت (با استفاده از چک‌لیست کسپ^۱)، سال انتشار (با تأکید بر منابع جدیدتر) و اعتبار منابع بودند. درنهایت، مقالات منتخب برای ورود به مرحله تحلیل انتخاب شدند. ج- استخراج و تحلیل داده‌ها: اطلاعات کلیدی از مقالات انتخاب شده استخراج و به‌صورت سیستماتیک کدگذاری و دسته‌بندی شدند. برای تحلیل داده‌ها از روش تحلیل مضمون استفاده شد. در این روش، ابتدا کدهای اولیه از متن مقالات استخراج شد. سپس، این کدها در قالب مضامین و زیرمضامین دسته‌بندی شدند. درنهایت، با تفسیر مضامین و زیر مضامین، به سؤال پژوهش، یعنی «چگونه می‌توان با استفاده از الگوریتم‌های یادگیری ماشین، اطلاعات نادرست و گمراه‌کننده را در جنگ‌شناختی شناسایی و خنثی کرد» پاسخ داده شد.

تحلیل داده‌ها در این پژوهش بر اساس مدل شش مرحله‌ای براون و کلارک (۲۰۰۶) صورت گرفته است. این فرآیند به‌صورت سیستماتیک در مراحل زیر اجرا شد:

۱. آشنایی با داده‌ها: مطالعه مکرر و عمیق متون مقالات منتخب جهت استغراق در داده‌ها
۲. ایجاد کدهای اولیه: استخراج کدهای اولیه (نظیر نشانگرهای زبانی شایعه، الگوهای ساختاری شبکه‌ها و لایه‌های پردازش عمیق) از متن مقالات.
۳. جستجوی مضامین: دسته‌بندی کدهای اولیه بر اساس قرابت مفهومی و شکل‌دهی به مضامین فرعی.
۴. بازبینی مضامین: کنترل و تطبیق مضامین استخراج‌شده با کل داده‌های متنی جهت اطمینان از روایی آن‌ها.
۵. تعریف و نام‌گذاری مضامین: ادغام، پالایش و نام‌گذاری نهایی مضامین که درنهایت ۵ مضمون اصلی (انواع اطلاعات نادرست، روش‌های شناسایی، الگوریتم‌های یادگیری، سطوح خنثی‌سازی و چالش‌های ساختاری) به‌عنوان شبکه مضامین پژوهش تبیین گردید.
۶. نگارش گزارش نهایی: تحلیل نهایی مضامین و پیوند آن‌ها با سؤالات پژوهش و ادبیات موجود.

تجزیه و تحلیل و یافته‌های پژوهش

در این بخش، یافته‌های حاصل از تحلیل کیفی مقالات انتخاب‌شده در مراحل مرور نظام‌مند، با استفاده از روش تحلیل مضمون ارائه می‌شود. ابتدا، با بررسی دقیق متون،

¹ CASP: Critical Appraisal Skills Programme

کدهای اولیه از مقالات استخراج شدند. سپس، با دسته‌بندی و ادغام کدهای مشابه، مضامین و زیر مضامین اصلی شکل گرفتند. درنهایت، با تفسیر و تحلیل مضامین و زیر مضامین، به سؤال پژوهش، یعنی «چگونه می‌توان با استفاده از الگوریتم‌های یادگیری ماشین، اطلاعات نادرست و گمراه‌کننده را در جنگ شناختی شناسایی و خنثی کرد؟» پاسخ داده شد.

مضامین و زیر مضامین

مضمون ۱: انواع اطلاعات نادرست در جنگ شناختی

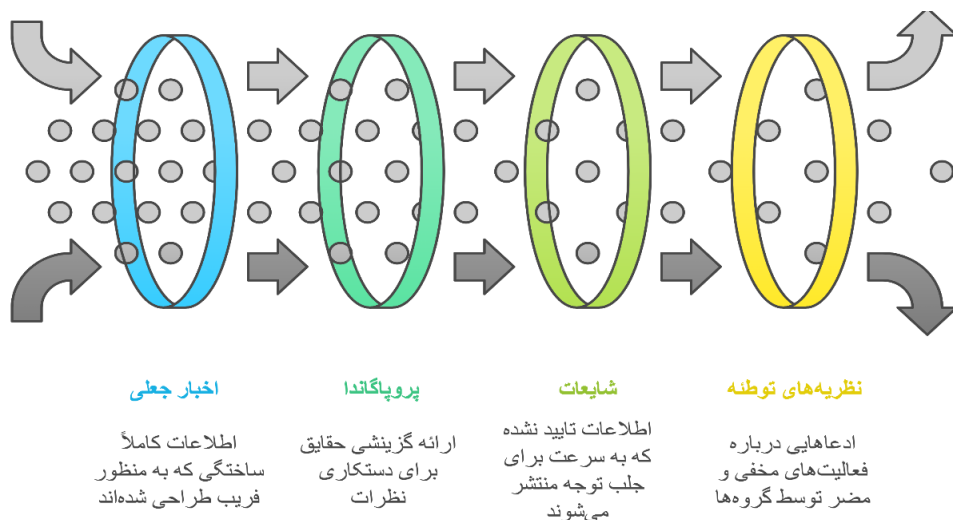
درک دقیق گونه‌شناسی اطلاعات نادرست، زیربنای طراحی هرگونه سیستم دفاعی مبتنی بر یادگیری ماشین است. در نبردهای شناختی، دشمن از طیف وسیعی از محتواهای تحریف‌شده برای نفوذ به لایه‌های ادراکی استفاده می‌کند که در این پژوهش به‌صورت زیر تبیین و بسط داده شده‌اند:

۱- اخبار جعلی و پروپاگاندا: اخبار جعلی و پروپاگاندا^۱ دو روی یک سکه‌اند که به‌طور هدفمند برای فریب، انحراف یا دست‌کاری افکار عمومی و فرآیندهای عالی ذهن در جنگ شناختی به کار می‌روند. اخبار جعلی شامل اطلاعاتی کاملاً ساختگی، کذب و فاقد مبنای واقعی هستند که با ظاهری مشابه اخبار معتبر و با استفاده از فن‌های روزنامه‌نگاری زرد، برای فریب مخاطب طراحی می‌شوند. در مقابل، پروپاگاندا لزوماً تماماً کذب نیست؛ بلکه ممکن است حاوی حقایقی باشد که به شکلی گزینشی، مبالغه‌آمیز و جهت‌دار برای پیشبرد اهداف سیاسی یا نظامی خاصی ارائه می‌شوند. این ترکیب «صدق و کذب» باعث می‌شود که پروپاگاندا نسبت به اخبار جعلی محض، قدرت نفوذ بیشتری در لایه‌های عمیق ذهنی فرماندهان داشته باشد و تشخیص آن توسط ماشین نیازمند تحلیل‌های زمینه‌ای و معنایی پیچیده‌تری باشد. هدف نهایی این ابزارها، تخریب آگاهی موقعیتی و جهت‌دهی به تصمیمات راهبردی نیروهای خودی است (البحر و البحر، ۲۰۲۰؛ هو و همکاران، ۲۰۲۴؛ علویطبی و الدوساری، ۲۰۲۴؛ سوئانتیرادوی و همکاران، ۲۰۲۱).

۲- شایعات و نظریه‌های توطئه: شایعات و نظریه‌های توطئه هر دو اطلاعاتی تأیید نشده و سیال هستند که به‌سرعت در فضای اطلاعاتی منتشر می‌شوند و می‌توانند باعث ایجاد بی‌ثباتی، ترس، ابهام و بی‌اعتمادی در جامعه و ساختار نظامی شوند. در مقابل، نظریه‌های

¹ Propaganda

توطئه ادعاهایی نظام‌مند هستند که معتقدند یک گروه، سازمان یا قدرت مخفی در حال انجام فعالیت‌های پنهانی، هماهنگ و مخرب برای کنترل رویدادها است. در نبردهای مدرن، این اطلاعات در جنگ‌شناختی به‌عنوان ابزاری قدرتمند برای تضعیف روحیه، ایجاد تفرقه در شبکه فرماندهی و کنترل و اختلال در فرآیندهای منطقی تصمیم‌گیری به کار می‌روند. این ابزارها با هدف قرار دادن سوگیری‌های شناختی فرماندهان، تاب‌آوری آن‌ها را در برابر داده‌های متناقض کاهش داده و منجر به فلج تصمیم‌گیری در لحظات حساس می‌شوند (هو و همکاران، ۲۰۲۴؛ سیربو^۱ و همکاران، ۲۰۲۳).



شکل (۱) انواع اطلاعات نادرست در جنگ‌شناختی (منبع: یافته‌های پژوهش حاضر)

مضمون ۲: روش‌های شناسایی اطلاعات نادرست با استفاده از الگوریتم‌های یادگیری ماشین

شناسایی به‌موقع و دقیق اطلاعات نادرست، اولین و حیاتی‌ترین گام در چرخه‌ی خنثی‌سازی تهدیدات در جنگ‌شناختی محسوب می‌شود. الگوریتم‌های یادگیری ماشین با تحلیل ابعاد مختلف داده‌ها، به فرماندهان کمک می‌کنند تا از میان توده‌ی عظیم اطلاعات آلوده، واقعیت را استخراج کرده و آگاهی موقعیتی خود را حفظ کنند. این روش‌ها در سه دسته‌ی اصلی زیر تبیین و بسط داده می‌شوند:

¹ Sîrbu et al

۱- روش‌های مبتنی بر محتوا: این روش‌ها با تحلیل مستقیم و عمیق محتوای اطلاعات (اعم از متن، تصویر یا ویدئو)، به شناسایی الگوهای زبانی و بصری منحصر به فردی می‌پردازند که اغلب در اطلاعات نادرست تعبیه می‌شوند. در تحلیل متنی، الگوریتم‌های پردازش زبان طبیعی^۱ با واکاوی ساختارهای نحوی و معنایی، قادر به تشخیص لحن جانب‌دارانه، بزرگنمایی، استفاده افراطی از کلمات احساسی و تهییج‌کننده هستند که با هدف دست‌کاری عواطف فرمانده و نیروها طراحی شده‌اند. همچنین در حوزه‌ی بصری، این الگوریتم‌ها می‌توانند دست‌کاری‌های پیکسلی و ناهماهنگی‌های موجود در تصاویر و ویدئوهای جعلی مانند جعل عمیق^۲ را که برای فریب ادراکی در صحنه نبرد استفاده می‌شوند، با دقت بالا ردیابی کنند. این تحلیل‌های محتوایی باعث می‌شود بار شناختی فرمانده در تفکیک گزارش‌های واقعی از جعلی به شدت کاهش یابد (البحر و البحر، ۲۰۲۰؛ هو و همکاران، ۲۰۲۴).

۲- روش‌های مبتنی بر شبکه: این روش‌ها به جای تمرکز بر پیام، بر «مسیر و نحوه انتشار» آن تمرکز دارند. با تحلیل شبکه انتشار در بسترهایی مانند شبکه‌های اجتماعی، الگوریتم‌ها الگوهای غیرطبیعی ارتباطات را که ویژگی بارز کمپین‌های جنگ شناختی است، شناسایی می‌کنند.

۳- روش‌های مبتنی بر رفتار کاربر: این رویکرد با تحلیل دقیق و مستمر رفتار کاربران در فضای مجازی و شبکه‌های اطلاعاتی، الگوهای کنشگری مرتبط با انتشار اطلاعات نادرست را ردیابی می‌کند. این الگوها شامل مواردی نظیر تواتر اشتراک‌گذاری، زمان‌بندی تعاملات و تاریخچه رفتاری کاربران و گروه‌ها است. الگوریتم‌های یادگیری ماشین با یادگیری رفتارهای عادی کاربران، می‌توانند انحرافات رفتاری (مانند فعالیت‌های ۲۴ ساعته یا انتشار محتوا از موقعیت‌های جغرافیایی متناقض) را شناسایی کنند. این تحلیل‌ها به سیستم‌های پشتیبانی تصمیم اجازه می‌دهد تا منابع خبری غیرقابل اعتماد را برچسب‌گذاری کرده و از نفوذ داده‌های مخرب به فرآیند تصمیم‌گیری فرمانده جلوگیری کنند که این امر مستقیماً منجر به افزایش ثبات و تاب‌آوری شناختی در شرایط بحران می‌گردد (سیربو و همکاران، ۲۰۲۳).

^۱ NLP

^۲ Deepfake



شکل (۲) روش‌های شناسایی اطلاعات نادرست با استفاده از الگوریتم‌های یادگیری ماشین (منبع: یافته‌های پژوهش حاضر)

مضمون ۳: الگوریتم‌های یادگیری ماشین برای شناسایی اطلاعات نادرست

انتخاب الگوریتم مناسب، هسته مرکزی یک سیستم پشتیبانی تصمیم ۱ برای مقابله با جنگ‌شناختی است. الگوریتم‌ها باید فراتر از یک فیلتر ساده عمل کرده و توانایی درک پیچیدگی‌های زبانی و بصری حملات دشمن را داشته باشند. در این پژوهش، سه دسته اصلی از این الگوریتم‌ها مورد واکاوی قرار گرفتند:

۱- یادگیری عمیق: یادگیری عمیق^۲ با بهره‌گیری از معماری‌های پیشرفته شبکه‌های عصبی مصنوعی (مانند CNN^۳ و RNN^۴) توانمندی خیره‌کننده‌ای در شناسایی الگوهای غیرخطی و بسیار پیچیده در داده‌های حجیم دارد. این الگوریتم‌ها با شبیه‌سازی فرآیندهای پردازش عصبی انسان، قادرند ویژگی‌های پنهان در انواع مختلف اطلاعات نادرست، از جمله اخبار جعلی، پروپاگاندا و به‌ویژه تصاویر و ویدئوهای دستکاری‌شده (جعل عمیق) را ردیابی کنند. در محیط جنگ‌شناختی که دشمن از چند رسانه‌ای‌های فریبده برای نفوذ به لایه‌های ناخودآگاه ذهن فرماندهان استفاده می‌کند، یادگیری عمیق با دقت

^۱ DSS

^۲ Deep Learning

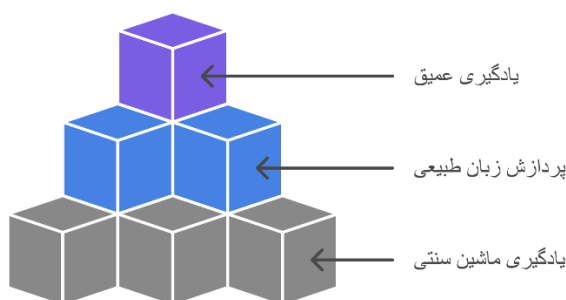
^۳ Convolutional Neural Network

^۴ Recurrent Neural Network

بالای خود، سپری ادراکی ایجاد کرده و از تخریب آگاهی موقعیتی جلوگیری می‌کند (البحر و البحر، ۲۰۲۰؛ سوئانتیرادوی و همکاران، ۲۰۲۱).

۲- پردازش زبان طبیعی: به‌عنوان پلی میان زبان انسانی و منطق ماشین، ابزاری حیاتی برای تحلیل محتوای متنی در جنگ‌شناختی است. این فناوری با تحلیل ساختارهای معنایی و نحوی، می‌تواند اطلاعات نادرست نهفته در اخبار جعلی و کمپین‌های پروپاگاندا را شناسایی کند. الگوریتم‌های پردازش زبان طبیعی با واکاوی مؤلفه‌هایی همچون تحلیل احساسات، تشخیص سوگیری و استخراج مضامین، به سیستم‌های پشتیبانی تصمیم اجازه می‌دهند تا مقاصد پنهان دشمن (مانند ایجاد هراس یا تضعیف اعتماد به نفس) را کشف نموده و فرمانده را نسبت به محتوای سمی آگاه سازند (علویطبی و الدوساری، ۲۰۲۴).

۳- یادگیری ماشین سنتی: الگوریتم‌های کلاسیک یادگیری ماشین نظیر درخت تصمیم^۱، ماشین بردار پشتیبان^۲ و بیز ساده^۳، به دلیل سرعت عملیاتی بالا و قابلیت تفسیرپذیری^۴، همچنان جایگاه ویژه‌ای در سامانه‌های نظامی دارند. شفافیت در تصمیم‌گیری این الگوریتم‌ها به فرمانده کمک می‌کند تا منطق پشت هشدار سیستم را درک کرده و با اطمینان بیشتری بر اساس خروجی سیستم پشتیبانی تصمیم عمل نماید (البحر و البحر، ۲۰۲۰).



شکل (۳) الگوریتم‌های یادگیری ماشین برای شناسایی اطلاعات نادرست
(منبع: یافته‌های پژوهش حاضر)

^۱ Decision Tree

^۲ SVM

^۳ Naive Bayes

^۴ Explainability

مضمون ۴: خنثی‌سازی اطلاعات نادرست در جنگ شناختی

شناسایی اطلاعات نادرست، تنها نیمی از مسیر ارتقاء تاب‌آوری شناختی است؛ نیمی دوم و حیاتی‌تر، خنثی‌سازی اثرات این اطلاعات بر ذهن فرمانده است. فرآیند خنثی‌سازی در این پژوهش در سه سطح عملیاتی تبیین شده است:

۱- حذف اطلاعات نادرست: در سطحی‌ترین و درعین‌حال سریع‌ترین لایه دفاعی، می‌توان از الگوریتم‌های یادگیری ماشین برای شناسایی و حذف خودکار اطلاعات مخرب از لایه‌های مختلف منابع انتشار و شبکه‌های ارتباطی نظامی استفاده کرد. این اقدام از ورود داده‌های آلوده به چرخه فرماندهی و کنترل جلوگیری کرده و مانع از اشباع شناختی فرمانده توسط گزارش‌های متناقض و دروغین می‌گردد (کلاوری و همکاران، ۲۰۲۱).

۲- برجسب‌گذاری اطلاعات نادرست: یکی از مؤثرترین روش‌ها برای تقویت تفکر انتقادی و تاب‌آوری شناختی، استفاده از سیستم‌های هوشمند برای شناسایی و برجسب‌گذاری خودکار محتوا است. این هشدارها به‌عنوان یک «سنگر شناختی» عمل کرده و باعث می‌شود فرمانده با دقت و تردید علمی بیشتری به تحلیل مطلب بپردازد که این خود مانع از پذیرش ناخودآگاه دروغ‌های دشمن می‌شود (هو و همکاران، ۲۰۲۴).

۳- ارائه اطلاعات صحیح: پیشرفته‌ترین لایه خنثی‌سازی، استفاده از الگوریتم‌های هوش مصنوعی برای پیشنهاد و ارائه هم‌زمان اطلاعات صحیح، مستند و معتبر در تقابل با اطلاعات نادرست است. این سیستم‌ها با شناسایی یک ادعای کذب، بلافاصله مستندات متناقض با آن‌را از پایگاه‌های داده معتبر استخراج کرده و به فرمانده ارائه می‌دهند. این تقابل داده‌ای^۱، باعث خنثی شدن اثر «تقدم» در پیام‌های فریبنده دشمن شده و بستر لازم برای اتخاذ تصمیمات مستقل و آگاهانه را در شرایط بحرانی فراهم می‌سازد (کرتیسووا، ۲۰۱۸).

^۱ Debunking



شکل (۴) خنثی‌سازی اطلاعات نادرست در جنگ شناختی

(منبع: یافته‌های پژوهش حاضر)

مضمون ۵: چالش‌ها و محدودیت‌های شناسایی و خنثی‌سازی اطلاعات نادرست علی‌رغم توانمندی‌های خیره‌کننده‌ی یادگیری ماشین، پیاده‌سازی این سیستم‌ها در محیط‌های حساس و پویای جنگ‌شناختی با چالش‌ها و محدودیت‌های بنیادینی روبرو است که نادیده گرفتن آن‌ها می‌تواند بر فرآیند تصمیم‌گیری راهبردی اثرات معکوس بگذارد. این چالش‌ها در محورهای زیر تبیین و تحلیل می‌شوند:

۱- کمبود و تنوع داده‌های آموزشی: برای آموزش دقیق و کارآمد الگوریتم‌های یادگیری ماشین، به مجموعه داده‌های آموزشی باکیفیت بالا، حجم بسیار زیاد و تنوع مناسب نیاز است. در حوزه‌های دفاعی، به دلیل محرمانگی و ماهیت نوظهور تهدیدات شناختی، دسترسی به چنین داده‌های جامعی به راحتی ممکن نیست. فقدان داده‌های کافی منجر به پدیده‌ی «بیش‌برازش»^۱ شده و قدرت تعمیم‌پذیری الگوریتم را در برابر حملات ناشناخته‌ی دشمن کاهش می‌دهد (احمد و همکاران، ۲۰۲۱؛ علویطبی و الدوساری، ۲۰۲۴).

۲- پیچیدگی و تغییرپذیری اطلاعات نادرست: در جنگ‌شناختی، دشمن از تاکتیک‌های «یادگیری ماشین خصمانه» استفاده می‌کند تا به‌طور مداوم روش‌های انتشار و

¹ Overfitting

تکنیک‌های تولید اطلاعات نادرست را تکامل دهد. این تغییرپذیری دائم، شناسایی را برای مدل‌های ایستا دشوار می‌سازد. ظهور فناوری‌هایی نظیر «جعل عمیق» در محتواهای ویدئویی و جعل صدا با استفاده از هوش مصنوعی^۱، تشخیص اطلاعات نادرست را به‌طور چشم‌گیری دشوار کرده است. این تکنیک‌ها با ایجاد شکاف بین «واقعیت» و «ادراک»، آگاهی موقعیتی فرمانده را هدف قرار داده و ماشین را در تشخیص مرزهای حقیقت با چالش‌های فنی جدی مواجه می‌کنند (هو و همکاران، ۲۰۲۴).

۳- محدودیت‌های ذاتی الگوریتم‌های یادگیری ماشین: الگوریتم‌ها باوجود قدرت محاسباتی، در شناسایی و خنثی‌سازی اطلاعات نادرست دارای محدودیت‌های ساختاری هستند. برای مثال، ماشین‌ها اغلب در درک مفاهیم انتزاعی، استعاره‌ها، زمینه‌های فرهنگی و معانی ضمنی^۲ ضعیف عمل می‌کنند و ممکن است به‌راحتی توسط تکنیک‌های فریبنده‌ی پیشرفته فریب داده شوند. علاوه براین، ریسک مثبت کاذب^۳ همواره وجود دارد. چنین خطایی در صحنه نبرد می‌تواند منجر به نادیده گرفتن گزارش‌های واقعی و حیاتی شده و اعتماد فرمانده به سیستم پشتیبانی تصمیم را خدشه‌دار کند (کرتیسووا، ۲۰۱۸؛ ساندو و همکاران، ۲۰۲۴).

۴- مسائل اخلاقی و حقوقی: استفاده از الگوریتم‌ها برای خنثی‌سازی اطلاعات می‌تواند با چالش‌های اخلاقی نظیر سانسور ناخواسته، نقض حریم خصوصی و سوگیری الگوریتمی^۴ همراه باشد. همچنین، اگر داده‌های آموزشی حاوی سوگیری‌های خاص (نژادی، مذهبی یا سیاسی) باشند، الگوریتم نیز دچار تبعیض شده و منجر به شناسایی نادرست اطلاعات یا حذف آرای منتقدان می‌شود. در یک محیط نظامی، این سوگیری‌ها می‌تواند «تونل‌زنی شناختی» ایجاد کرده و مانع از درک همه‌جانبه‌ی واقعیت توسط فرمانده گردد (میلر، ۲۰۲۳؛ پاولیکی و همکاران، ۲۰۲۴).

۵- ضرورت تعامل انسان و ماشین (انسان در حلقه): شناسایی و خنثی‌سازی اطلاعات نادرست نباید به‌طور کاملاً خودمختار به ماشین واگذار شود. الگوریتم‌های یادگیری ماشین باید به‌عنوان ابزارهای کمکی در خدمت «ارتقاء آگاهی» انسان باشند. اگرچه ماشین در

^۱ Voice Cloning

^۲ Semantics

^۳ False Positive

^۴ Algorithmic Bias

پردازش حجم عظیم داده‌ها بی‌رقیب است، اما قضاوت نهایی در مورد صحت اطلاعات، تحلیل پیامدهای سیاسی-نظامی و اتخاذ تصمیمات اخلاقی و حساس باید بر عهده فرمانده (انسان) باشد. تقویت این تعامل، کلید اصلی دستیابی به تاب‌آوری شناختی است؛ چراکه تکیه‌ی مطلق به هوش مصنوعی می‌تواند منجر به «تنبلی شناختی» و کاهش قدرت تحلیل مستقل فرمانده در شرایط بحران شود (ون دن بوش و برونکه‌ورست، ۲۰۱۸؛ شوبرت و همکاران، ۲۰۱۸).



شکل (۵) چالش‌ها و محدودیت‌های شناسایی و خنثی‌سازی اطلاعات نادرست
(منبع: یافته‌های پژوهش حاضر)

بحث و نتیجه گیری

این پژوهش با استفاده از روش مرور نظام مند و تحلیل مضمون، به بررسی این سؤال پرداخت که چگونه می توان با استفاده از الگوریتم های یادگیری ماشین، اطلاعات نادرست و گمراه کننده را در جنگ شناختی شناسایی و خنثی کرد؟

یافته های این پژوهش نشان داد که الگوریتم های یادگیری ماشین، با بهره گیری از روش های مختلف شناسایی مانند تحلیل محتوا، شبکه و رفتار کاربر، می توانند انواع مختلف اطلاعات نادرست، از جمله اخبار جعلی، پروپاگاندا، شایعات و نظریه های توطئه را شناسایی کنند.

همچنین، مشخص شد که الگوریتم های یادگیری عمیق، پردازش زبان طبیعی و یادگیری ماشین سنتی می توانند در شناسایی اطلاعات نادرست مؤثر باشند. برای خنثی سازی اطلاعات نادرست، روش هایی مانند حذف، برچسب گذاری، ارائه اطلاعات صحیح و آموزش و افزایش آگاهی مورد بررسی قرار گرفتند.

با این حال، چالش ها و محدودیت هایی مانند کمبود داده های آموزشی با کیفیت، حجم زیاد و تنوع مناسب، پیچیدگی و تغییرپذیری اطلاعات نادرست، محدودیت های الگوریتم های یادگیری ماشین، مسائل اخلاقی و نیاز به تعامل انسان و ماشین نیز در این زمینه وجود دارند.

به طور کلی، این پژوهش نشان داد که استفاده از الگوریتم های یادگیری ماشین در کنار سایر روش ها و با توجه به چالش های موجود، می تواند به طور مؤثر در شناسایی و خنثی سازی اطلاعات نادرست در جنگ شناختی مورد استفاده قرار گیرد.

اهمیت این یافته ها در زمینه جنگ شناختی بسیار قابل توجه است، زیرا شناسایی و خنثی سازی به موقع اطلاعات نادرست می تواند از ایجاد اختلال در فرآیند تصمیم گیری فرماندهان، کاهش روحیه نیروها و ایجاد شکاف بین مردم و نیروهای نظامی جلوگیری کند.

در نهایت، با توجه به اهمیت موضوع و یافته های این پژوهش، پیشنهاد می شود که در آینده تحقیقات بیشتری در زمینه شناسایی و خنثی سازی اطلاعات نادرست در جنگ شناختی با استفاده از الگوریتم های یادگیری ماشین انجام شود.

تبیین اهمیت یافته‌ها

یافته‌های این پژوهش می‌تواند برای فرماندهان نظامی، طراحان سیستم‌های هوش مصنوعی و پژوهشگران حوزه جنگ شناختی مفید باشد.

فرماندهان نظامی؛ با آگاهی از قابلیت‌های الگوریتم‌های یادگیری ماشین در شناسایی اطلاعات نادرست، می‌توانند از این الگوریتم‌ها به‌طور مؤثرتری برای مقابله با تهدیدات جنگ شناختی استفاده کنند.

طراحان سیستم‌های هوش مصنوعی؛ با در نظر گرفتن چالش‌ها و محدودیت‌های شناسایی اطلاعات نادرست، می‌توانند سیستم‌های هوش مصنوعی کارآمدتر و قابل‌اعتمادتری برای این منظور طراحی کنند.

پژوهشگران؛ این پژوهش می‌تواند زمینه‌ساز تحقیقات بیشتر در حوزه شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ شناختی باشد.

پیشنهادهای

تحقیقات آینده

- انجام مطالعات کمی و تجربی برای تأیید یافته‌های این پژوهش و بررسی اثربخشی الگوریتم‌های یادگیری ماشین در شناسایی اطلاعات نادرست در محیط‌های واقعی.
- بررسی تأثیر عوامل فردی، سازمانی و فرهنگی بر اثربخشی الگوریتم‌های یادگیری ماشین در شناسایی اطلاعات نادرست.
- توسعه الگوریتم‌های یادگیری ماشین جدید برای شناسایی انواع مختلف اطلاعات نادرست، از جمله دیپ فیک‌ها و جعل صدا.
- توسعه روش‌های ترکیبی برای شناسایی و خنثی‌سازی اطلاعات نادرست که از نقاط قوت الگوریتم‌های مختلف یادگیری ماشین و نیروی انسانی بهره می‌برند.

کاربرد یافته‌ها در عمل

- آموزش فرماندهان نظامی در مورد قابلیت‌ها و چالش‌های الگوریتم‌های یادگیری ماشین در شناسایی اطلاعات نادرست.
- طراحی و پیاده‌سازی سیستم‌های هوش مصنوعی برای شناسایی و خنثی‌سازی اطلاعات نادرست در سازمان‌های نظامی.

- آموزش و افزایش آگاهی در جامعه در مورد چگونگی شناسایی اطلاعات نادرست.

محدودیت‌ها

- تعداد مطالعاتی که به‌طور خاص به موضوع شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ شناختی با استفاده از الگوریتم‌های یادگیری ماشین پرداخته‌اند، محدود است.
- بیشتر مطالعات انجام شده در این حوزه، از نوع کیفی بوده‌اند و نیاز به انجام مطالعات کمی و تجربی برای تأیید یافته‌ها وجود دارد.
- باوجود این محدودیت‌ها، این پژوهش می‌تواند به‌عنوان یک گام اولیه در بررسی نقش الگوریتم‌های یادگیری ماشین در شناسایی و خنثی‌سازی اطلاعات نادرست در جنگ شناختی مفید باشد.

منابع

منابع فارسی

- افشردی، محمد؛ و شیخ، علیرضا. (۱۳۹۸). بررسی نیازها، قابلیت‌ها و توانمندی‌های فرماندهی و مدیریت در عرصه جنگ‌های آینده ج.ا.ایران. راهبرد دفاعی، ۱۷(۲)، ۴۹-۶۸.

(URL: https://ds.sndu.ac.ir/article_806.html)

- بختیاری، ایرج؛ و عسکری، احمد. (۱۴۰۱). ارائه مدل آگاهی وضعیتی در شبکه فرماندهی و کنترل پدافند هوایی. مطالعات دفاعی استراتژیک، ۲۰(۸۸)، ۱۲۳-۱۵۰.

(URL: https://sds.sndu.ac.ir/article_1832.html)

- ترسلی، ا. (۱۴۰۳). کاربرد هوش مصنوعی در تشخیص اخبار جعلی. فصلنامه آماد و فناوری دفاعی، ۷(۲)، ۸۵-۱۱۲.

(URL: https://amfad.sndu.ac.ir/article_3047.html)

- رضایی، محسن؛ رشید، غلامعلی؛ و پوردستان، احمدرضا. (۱۳۹۹). مؤلفه‌ها و ویژگی‌های فرماندهی و کنترل هوشمند در صحنه نبرد. علوم و فنون نظامی، ۱۶(۵۴)، ۱۴۹-۱۷۱.

(URL: https://qjmst.sndu.ac.ir/article_1028.html)

- شاملو، رضا. (۱۴۰۱). جایگاه فناوری هوش مصنوعی در سامانه‌های فرماندهی و کنترل آینده (مطالعه موردی: سامانه فرماندهی و کنترل فراگیر مشترک آمریکا). نشریه مطالعات جنگ، ۴ (۱۵)، ۸۱-۱۰۶.

(URL: <https://jws.ir/Article-1-1118-fa.html>)

- فرهنگ، س؛ و آروند، ح. (۱۴۰۲). کاربرد هوش مصنوعی در بازی جنگ برای بهبود عملکرد شناختی در نبردهای دریایی ترکیبی. مدیریت نظامی، ۲۳ (۹۲)، ۷۸-۵۷.

(URL: <https://qjmm.ir/article-1-2401-fa.html>)

- کنعانی، ا؛ و یآوری، ع. (۱۴۰۱). طراحی الگوی فرماندهی و کنترل انتظامی هوشمند. پژوهش‌های مدیریت انتظامی (مطالعات مدیریت انتظامی)، ۱۷ (۲)، ۶۵-۹۱.
- محمودی، ابراهیم؛ و بحرکاظمی، مهران. (۱۴۰۳). هوش مصنوعی و تأثیر آن بر امنیت سایبری و حفاظت از داده‌ها. فصلنامه آماد و فناوری دفاعی، ۷ (۱)، ۴۵-۶۸.
- مداح، پ؛ و شاه‌محمدی، م. (۱۴۰۲). فرصت‌ها و چالش‌های هوش مصنوعی در جنگ‌شناختی. دانشنامه علوم سیاسی، ۴ (۱۳)، ۸۴-۱۰۱.

(URL: https://psj.sndu.ac.ir/article_2634.html)

منابع انگلیسی

- Afshordi, M. & Sheikh, A. (2019). Examining the needs, capacities and capabilities of command and management in the field of future wars of the I.R. Iran. Strategic Defense, 17(2), 49-68. [In Persian] (URL: https://ds.sndu.ac.ir/article_806.html)
- Ahmed, A. A. A. Aljabouh, A. Donepudi, P. K. & Choi, M. S. (2021). Detecting fake news using machine learning: A systematic literature review. arXiv preprint arXiv:2102.04458.
- Albahr, A. & Albahr, M. (2020). An empirical comparison of fake news detection using different machine learning algorithms. International Journal of Advanced Computer Science and Applications, 11(9).
- Alotaibi, T. & Al-Dossari, H. (2024). A Review of Fake News Detection Techniques for Arabic Language. International Journal of Advanced Computer Science & Applications, 15(1).
- Bakhtiari, I. & Askari, A. (2022). Presenting a situational awareness model in the air defense command and control network. Strategic Defense

- Studies, 20(88), 123-150. [In Persian] (URL: https://sds.sndu.ac.ir/article_1832.html)
- Buvarp, P. M. H. (2023). Technical Evaluation Report (TER) HFM-361 Research Symposium (RSY) Mitigating and Responding to Cognitive Warfare.
 - Claverie, B. & Du Cluzel, F. (2022). "Cognitive warfare": The advent of the concept of "cognitics" in the field of warfare. *Cognitive Warfare: the future of cognitive dominance*, 2(2), 1-7.
 - Claverie, B. Prébot, B. Buchler, N. & Du Cluzel, F. (2021). *Cognitive Warfare: The Future of Cognitive Dominance*. Bruxelles: NATO Science and Technology Organization.
 - Farhang, S. & Arvand, H. (2023). Application of artificial intelligence in wargaming to improve cognitive performance in hybrid naval battles. *Military Management*, 23(92), 57-78. [In Persian]
 - Fenstermacher, L. Uzcha, D. Larson, K. Vitiello, C. & Shellman, S. (2023). New perspectives on cognitive warfare. *Signal Processing, Sensor/Information Fusion, and Target Recognition XXXII*.
 - Greipl, A. Davison, N. & Hinds, G. (2024). Artificial Intelligence and Related Technologies in Military Decision-Making on the Use of Force in Armed Conflicts. (URL: <https://www.geneva-academy.ch>)
 - Hu, B. Mao, Z. & Zhang, Y. (2024). An overview of fake news detection: From a new perspective. *Fundamental Research*. <https://doi.org/10.1016/j.fmre.2024.01.017>
 - Kanaani, A. & Yavari, A. (2022). Designing a model for command and control in intelligent policing. *Journal of Police Management Studies*, 17(2), 65-91. [In Persian]
 - Kertysova, K. (2018). Artificial intelligence and disinformation: How AI changes the way disinformation is produced, disseminated, and can be countered. *Security and Human Rights*, 29, 55-81. <https://doi.org/10.1163/18750230-02901005>
 - Kostopoulos, G. Davrazos, G. & Kotsiantis, S. (2024). Explainable Artificial Intelligence-Based Decision Support Systems: A Recent Review. *Electronics*, 13(14), 2842. <https://doi.org/10.3390/electronics13142842>
 - Le Guyader, H. (2022). Cognitive domain: A sixth domain of operations. *Cognitive Warfare: the future of cognitive dominance*, 3, 1-5.
 - Maddah, P. & Shahmohammadi, M. (2023). Opportunities and challenges of artificial intelligence in cognitive warfare. *Encyclopedia of Political Science*, 4(13), 84-101. [In Persian]
 - Mahmoudi, A. & Bahrkazemi, M. (2024). Artificial intelligence and its impact on cybersecurity and data protection. *Amad and Defense Technology Journal*, 7(1), 45-68. [In Persian]
 - Miller, S. (2023). Cognitive warfare: an ethical analysis. *Ethics and Information Technology*, 25(3), 46.

- Morgan, F. E. et al. (2020). Military applications of artificial intelligence. Santa Monica: RAND Corporation.
- Nadibaidze, A. Bode, I. & Zhang, Q. (2024). AI in Military Decision Support Systems: A Review of Developments and Debates.
- P, S. (2020). Detection of Propaganda from News Articles using Deep Learning. International Journal of Advanced Trends in Computer Science and Engineering, 9, 3500-3505. <https://doi.org/10.30534/ijatcse/2020/155932020>
- Pawlicki, M. et al. (2024). Advanced insights through systematic analysis: Mapping future research directions and opportunities for xAI in deep learning and AI used in cybersecurity. Neurocomputing, 590, 127759.
- Rezaei, M. Rashid, G. & Pourdastan, A. (2020). Components and characteristics of intelligent command and control in the battlefield. Military Science and Tactics, 16(54), 149-171. [In Persian]
- Sandu, A. et al. (2024). Machine Learning and Deep Learning Applications in Disinformation Detection: A Bibliometric Assessment. Electronics, 13(22), 4352.
- Schubert, J. et al. (2018). Artificial intelligence for decision support in command and control systems. Proceedings of the 23rd ICCRTS.
- Sedova, K. (2021). AI and the Future of Disinformation Campaigns: Part 2, a Threat Model. CSET.
- Shamloo, R. (2022). The position of artificial intelligence technology in future command and control systems (Case study: US Joint All-Domain Command and Control System). War Studies Journal, 4(15), 81-106. [In Persian]
- Tarsali, A. (2024). Application of artificial intelligence in fake news detection. Amad and Defense Technology, 7(2), 85-112. [In Persian]
- Van Den Bosch, K. & Bronkhorst, A. (2018). Human-AI cooperation to benefit military decision making.
- Wang, S. Bao, Y. & Li, Y. (2018). The architecture and technology of cognitive electronic warfare. Sci. Sin. Inf, 48(12), 1603-1613.
- Xingdong, F. & Xiangming, Z. (2022). Algorithmic Cognitive Warfare: Paradigm Shift of Public Opinion Warfare in the Context of Russia-Ukraine Conflict. Media Observer, 460(4), 5-15.