

Dazzle 2.0

Russia's “zebra” trucks and the contest between a can of paint and a machine that thinks it can see.

What the viral claim gets right, what it inflates, and what it means for how the West buys autonomy. The first paper in the NPSI Counter-Autonomy series.

AUTHOR	Jesse James
DOCUMENT ID	NPSI-WP-007
VERSION	v1.0 · change log
SERIES	NPSI Counter-Autonomy — the contest between machine autonomy and its countermeasures
RELEASED	12 July 2026 · Victoria, British Columbia
STATUS	Unclassified · All sources open. Editorially controlled. Public version-controlled review open.
CITE AS	James, J. (2026). <i>Dazzle 2.0</i> (NPSI Working Paper No. 7, v1.0). https://npsi.ca/wp7/
COMPANIONS	TB1 — The Verified Sky (the sensing substrate) WP4 — The Addition Paradox (the strategic preface)

EXECUTIVE SUMMARY

The viral claim is **real at its root and inflated at its tip**. Russian logistics trucks in Ukraine have genuinely worn high-contrast black-and-white “dazzle” paint since late May 2026, and the intent — degrading the machine vision on Ukrainian strike drones — is accepted by serious analysts. But the social-media post that carried this claim laundered a *hedged* Economist feature into a confident assertion that the tactic *works*. No controlled test shows it defeating any named detector, and most analysts doubt it survives a thermal sensor or a human in the loop.

The engineering premise is legitimate; the crude execution probably is not. Peer-reviewed work proves physical patterns *can* fool object detectors — but the reliable attacks are mathematically optimized against a specific model, not hand-painted stripes. The stripes most plausibly work by pushing a truck *out of the training distribution*, an effect that evaporates the moment the detector is retrained on the new pattern.

Bottom line: the durable finding is not “paint beats AI.” It is *cost asymmetry*: a paint can costs nothing and iterates in an afternoon; a fielded detector re-learns only as fast as its slowest update loop — which a serving U.S. commander has put at *six months*. That asymmetry, not the zebra pattern itself, is the procurement warning for NATO and Canada: vision-only terminal guidance is brittle, sensor fusion and human authorization are the robust answers, and Western munitions are exposed to exactly the same trick.

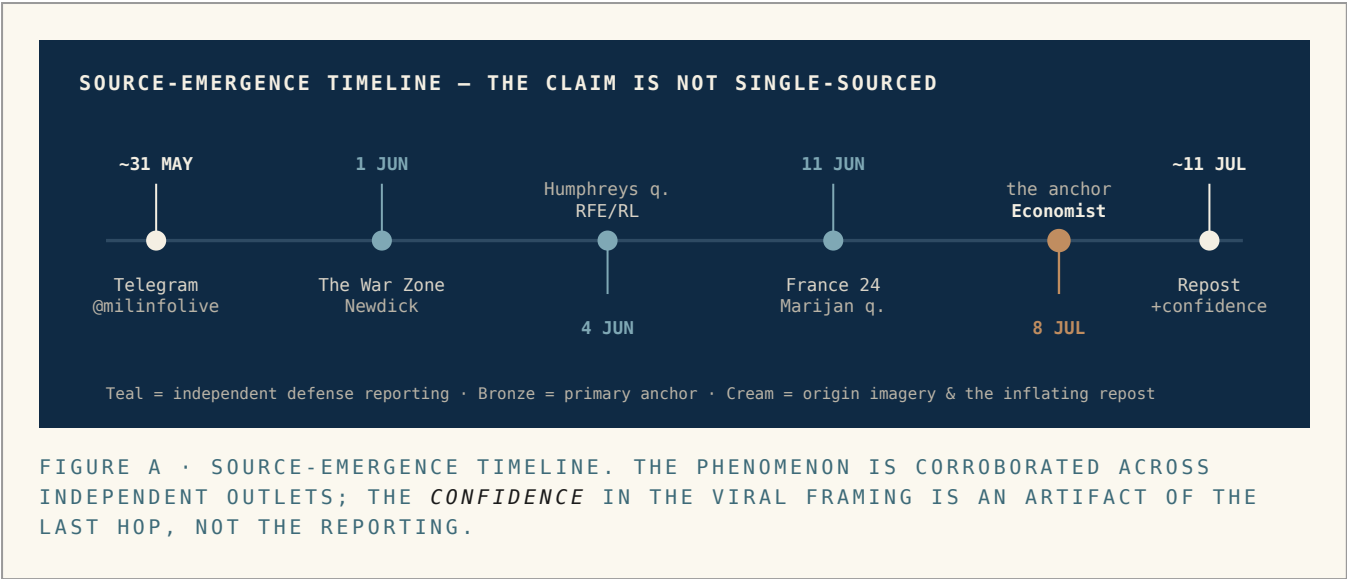
1. WHAT THE CLAIM ACTUALLY SAYS

The claim is well-corroborated in its existence and badly overstated in its effect — and the gap between those two things was manufactured by the repost chain, not by the reporting.

The image that started this file is a social-media post: a KamAZ truck painted like a zebra, over text asserting that Russia is confusing “AI-powered vision systems” and “causing drones to misidentify their targets,” sourced to *The Economist*. That sentence does something the underlying reporting never did: it states, flatly, that the tactic works.

The anchor is *The Economist*’s “How to hide from killer drones” (Science & technology, 8 July 2026), whose own standfirst frames the phenomenon as “bizarre forms of camouflage” in an emerging arms race — not as a Russian success. The article predicts that any advantage is *temporary*, because once zebra trucks are common enough they enter the training data and the models learn them. The repost keeps that final hedge but deletes the conditionality around everything before it. “Causing drones to misidentify their targets” is an outcome no cited source demonstrates.

This matters for a hostile reader, because the strongest attack on the whole story is exactly this inflation. So the first task is to separate the claim’s spine from its marketing. Three facts frame it: at least six independent outlets reported the trucks before or alongside the Economist anchor; the first imagery surfaced on Russian and Ukrainian Telegram channels around 31 May 2026; and the number of published controlled tests showing an effect on a named drone detector is zero.



2. CONFIRMED, PLAUSIBLE, OR HYPE

If a claim's tier cannot be stated, the claim cannot be defended — so every load-bearing statement in this paper is tagged.

A hostile analyst will not attack the trucks; the trucks are photographed. They will attack the leap from “trucks exist” to “trucks defeat AI.” The table below is the spine of the paper's credibility: it concedes everything that should be conceded, which is what earns the right to the findings that survive.

TIER	CLAIM	BASIS / CAVEAT
CONFIRMED	KamAZ and Ural trucks wear two high-contrast schemes (straight zebra; organic swirl), white over dark green, covering body, cab, wheels and tires.	Imagery across The War Zone, RFE/RL, France 24, Defense Express, Militarnyi since ~31 May 2026.
CONFIRMED	Physical patterns <i>can</i> defeat deep-learning vision in the real world.	Eykholt et al. (CVPR 2018): black/white stickers → 84.8% misclassification in a moving-vehicle field test.
CONFIRMED	The analogous 2023 aircraft-tire tactic degraded U.S. computer-vision identification.	Schuyler Moore, then CENTCOM CTO, on record at CSIS, September 2024.
CONFIRMED	Thermal/IR and a human in the loop defeat visible-spectrum paint.	Near-universal analyst consensus; Brave1 states humans always authorize strikes.

TIER	CLAIM	BASIS / CAVEAT
PLAUSIBLE	The trucks' <i>purpose</i> is specifically anti-AI.	Analyst and mil-blogger inference. No Russian MoD statement.
PLAUSIBLE	The paint works by out-of-distribution confusion, not optimized attack.	Theoretically sound (Humphreys); not field-measured on these trucks.
SPECULATION	The paint measurably degrades a real Ukrainian drone's targeting today.	No controlled test; no strike telemetry published. This is the inflated tip.
HYPE	Hand-painted stripes act as engineered adversarial examples or a durable counter.	Far more likely fragile out-of-distribution confusion; garish paint makes trucks <i>more</i> visible to eyes and thermal.

3. THE ENGINEERING

The literature says a detector can be fooled with paint; it also says it generally cannot be fooled with paint that was not computed.

How the drone actually sees

An AI-enabled strike drone runs a deep object detector — the YOLO family is the archetype for real-time inference on a small edge processor — that emits bounding boxes, class labels and confidence scores. Terminal guidance then locks onto a chosen detection and holds the aim point through the dive using optical-flow tracking. A human-in-the-loop system inserts an operator to confirm the target before or during that terminal phase. Each of those stages is a different thing to attack, and the paint plausibly bites at only one of them.

The proof that paint can attack vision

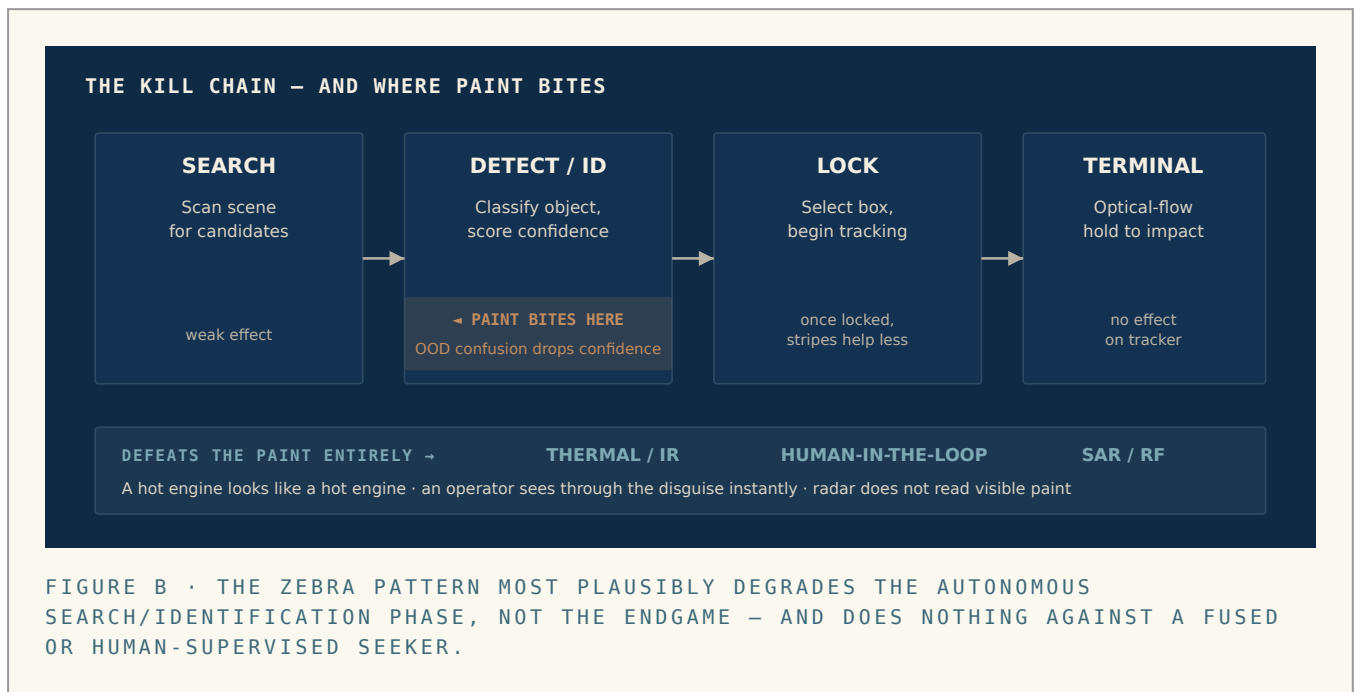
The adversarial-ML record is not ambiguous. Brown et al. (2017) built universal printable *adversarial patches* that force a chosen class regardless of scene. Eykholt et al. (2018) put black-and-white stickers on a real stop sign and caused targeted misclassification in 100% of lab images and 84.8% of frames from a moving vehicle. Thys, Van Ranst and Goedemé (2019) printed a patch that rendered a person invisible to a YOLOv2 detector. Vehicle-scale work — CAMOU, DAS, FCA — optimizes a texture across the whole three-dimensional surface of a car so it fails detection from many angles. The ceiling is real.

Why the Russian floor is far below that ceiling

Every one of those attacks was *computed* against a known model with gradient access. Hand-painted zebra stripes are not. Three gaps separate the two:

- **The optimization gap.** Effective patches are the solution to an optimization problem; stripes are folk art. Naive patterns consistently underperform optimized ones in the literature.
- **The physical-robustness gap.** Even optimized patches decay under lighting, viewpoint, distance, motion blur and paint fidelity — and a drone diving from altitude sweeps through all of those at once.
- **The transferability gap.** Attacks crafted on one architecture transfer unreliably to others, and especially weakly from convolutional networks to transformer detectors. Ukraine fields heterogeneous, frequently-updated models, so one pattern is unlikely to generalize.

The honest description of what the stripes most likely do is Todd Humphreys' own: they push the vehicle *out of the distribution* of the training images, so the model “might not know what it is looking at.” That is a genuine effect — and a fragile one. It collapses the moment zebra trucks enter the training set. Which is precisely why the Economist, and every serious analyst, calls the advantage temporary.

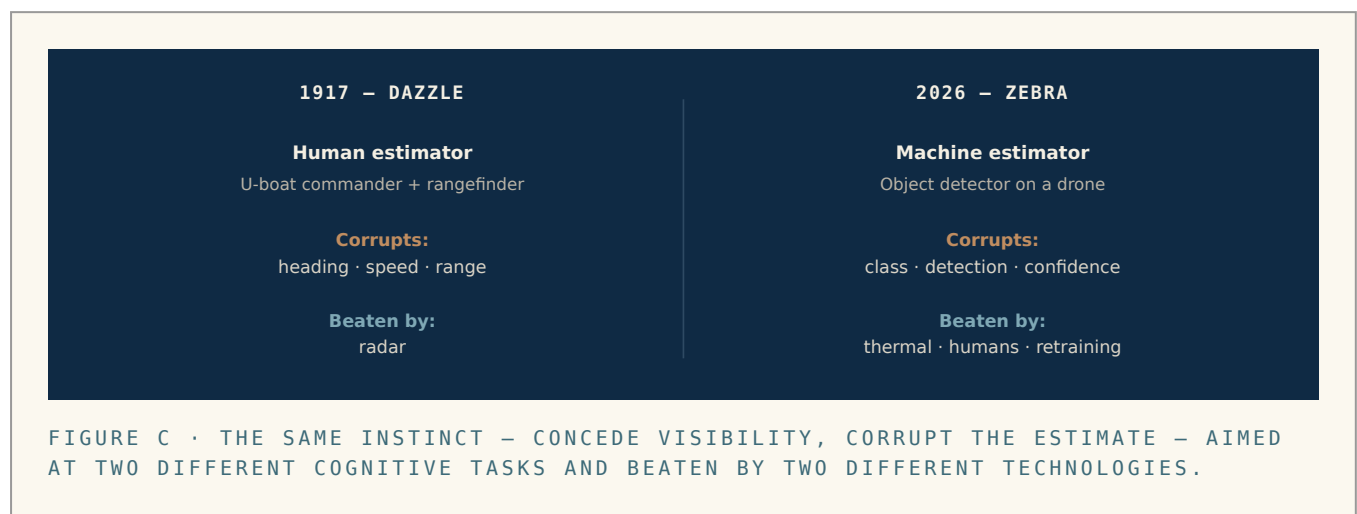


4. THE HISTORICAL HOOK

Dazzle targeted a human estimator; this targets a machine one. The century-old analogy is elegant and imperfect — an analogy, never an equivalence.

In 1917 the marine artist Norman Wilkinson devised what the Royal Navy called *dazzle*: violent geometric paint on merchant hulls that made no attempt to hide the ship. A smoke-belching freighter cannot hide on open ocean. The point was to corrupt the *human* optical estimate of the ship's course, speed and range as a U-boat commander read it through a periscope rangefinder, so the torpedo went where the ship wasn't. Britain painted thousands of vessels; the United States over a thousand. Whether it ever measurably worked was contested then and is contested now — the most rigorous recent study (Lovell, Sharman and Meese, 2024) found only a small perceptual twist of roughly ten degrees.

The 2026 version swaps the estimator. A machine, not a submariner, now judges what the object is — and the paint tries to corrupt that judgment instead. That is a clean and genuinely useful frame. But it is loose in one way worth conceding up front: First World War dazzle attacked *geometric estimation* (heading and range); the zebra truck attacks *object classification*. Defense Express's critique — that Russia has borrowed the aesthetic while missing the mechanism — is fair on the physics, even if the broader anti-recognition intent is coherent.



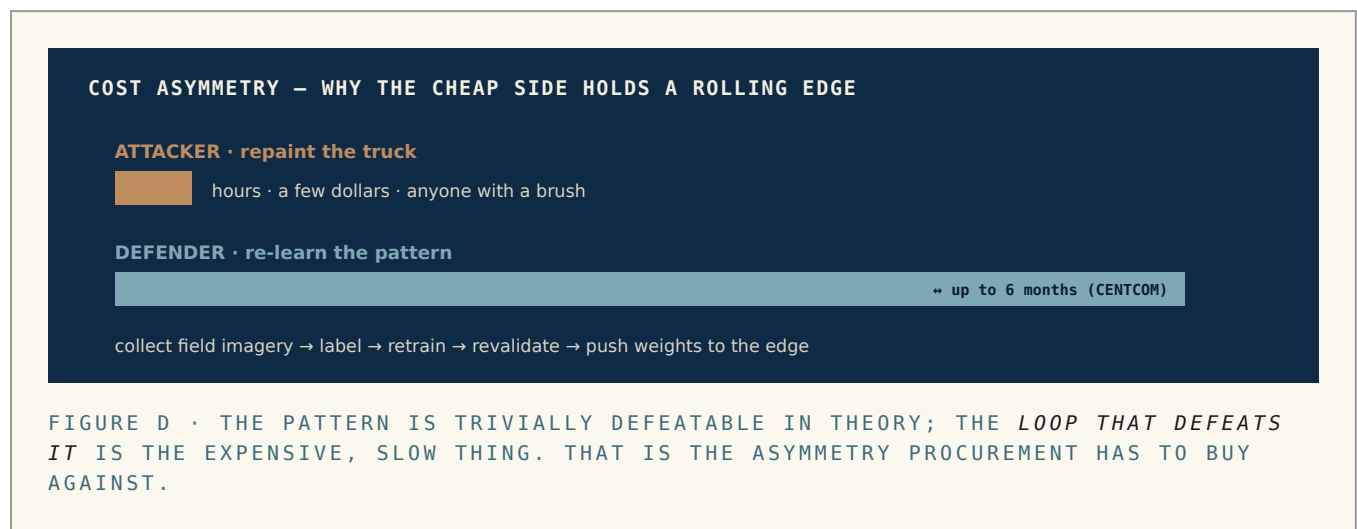
5. THE REAL STORY — COST ASYMMETRY

“AI adapts, so any advantage is temporary” is directionally right and strategically incomplete — because it assumes the adaptation loop is fast, and a U.S. commander says it is not.

The best-sourced evidence in this whole file is not about trucks. From around August 2023, Russia parked strategic bombers with car tires laid on their wings. In September 2024, Schuyler Moore — CENTCOM's first Chief Technology Officer — confirmed on the record that the trick works: put tires on a wing and many computer-vision models stop recognizing the aircraft as a plane. But her

operational point was sharper than the trick itself. If it takes six months to retrain the model, the adversary simply changes the pattern the next day and breaks it again — “we’re spending inordinate amounts of time on computer vision with very little to gain.” Her prescription was to push labeling and retraining to the tactical edge so the loop closes in days.

That reframes the viral hedge. Yes, the model eventually learns the zebra. But “eventually” is a variable, and the cost of resetting it is a can of paint. Repainting iterates faster and cheaper than a field machine-learning pipeline in any regime where the seeker is vision-only, un-fused, and slowly updated. The correct statement is not “AI wins” but: *AI wins wherever fusion, human oversight, and fast edge-retraining are present; cheap paint can win transiently wherever they are absent.*



6. DOCTRINE AND PROCUREMENT — FOR NATO, AND SPECIFICALLY FOR CANADA

The episode is not a Russian breakthrough; it is a cheap stress-test that Western autonomy would fail in exactly the same way unless it is bought correctly.

Ukrainian operators are publicly unbothered — a regiment commander vowed to burn the trucks “whatever they paint themselves as,” and Brave1 insists a human authorizes every strike. The pressure driving the Russian experiment is Ukraine’s mid-range strike campaign against logistics; the point of the paint is to buy a truck a few more kilometres of road. The strategic lesson travels the other way.

Canada is scaling loitering munitions and uncrewed systems quickly — Switchblade 300/600 to the Latvia brigade, the Army’s Minerva Initiative for uncrewed systems, and the CALM and CADUC deep-strike concepts. Every one of those programs will field seekers that a thirty-dollar

paint job is designed to spoof. The buy-side implications are concrete:

1. **Require multi-sensor seekers, not vision-only.** Mandate electro-optical plus thermal/IR fusion at acceptance. Visible-spectrum paint is defeated by heat; a vision-only munition is defeated by a brush.
2. **Mandate human-on-the-loop authorization** for target confirmation. It is the single most robust counter to physical deception and to the international-humanitarian-law risk of a detector over-trained to strike a pattern.
3. **Demand edge-retrainable models and an organic labeling pipeline.** Buy Moore's prescription directly: an update loop measured in days, not the six-month cycle that makes cheap paint a winning bet.
4. **Test automatic target recognition against adversarial and out-of-distribution inputs at acceptance.** Dazzle, decoys, tire-tactics, and AI-designed textures belong in the acceptance suite, not in the after-action report.
5. **Treat any single-vendor, black-box, slowly-updated recognition system as a liability,** not a feature. A model that cannot be retrained cedes the tempo to the adversary.

What would change this assessment

Three signals would move the zebra story from folk art to engineered threat: (1) a controlled test quantifying a confidence drop on a *named* detector; (2) verified strike footage showing lock-break attributable to paint; (3) either side shifting to thermal-only or fused seekers, or to *AI-designed* adversarial vehicle textures rather than hand-painted stripes. Watch for the third — it is the tell that the arms race has matured.

7. GUARDRAILS

The paper leads with the findings that survive a hostile read, and names the weak links before an opponent can.

Hostile-read weak links, stated plainly

The entire chain rests on social-media imagery plus expert inference; no quantitative field evidence exists. The expert most bullish on the tactic (Humphreys) still calls it short-lived. The most quotable skeptics (Defense Express; Wes O'Donnell) argue the mechanism is misunderstood. Individual photos are largely un-geolocated — treat them as illustrative, and stay alert to recycled or misattributed images, though none here are confirmed fake.

The Economist article is paywalled; its wording here is reconstructed from a verbatim republication and a faithful aggregator paraphrase, cross-confirmed by a third outlet. The Hornet munition's degree of autonomy is contested and its unit cost is reported across a wide

range — and a human in the loop nullifies the paint regardless. Adversarial-ML results are lab- and simulation-heavy; their transfer to a diving munition against a moving, repainted truck under battlefield lighting is not established and should not be assumed.

8. SELECTED SOURCES

- **The Economist.** “How to hide from killer drones.” Science & technology, 8 July 2026. *Primary anchor; paywalled; unsigned per house style.*
 - **Newdick, Thomas.** “Russian Trucks Get 'Dazzle' Paint To Throw Off AI-Enabled Drones.” The War Zone, 1 June 2026. twz.com <<https://www.twz.com/land/russian-trucks-get-dazzle-paint-to-throw-off-ai-enabled-drone>>
 - **Chapple, Amos.** “Zebra Trucks: Why Russia Is Using 'Dazzle Camouflage' In Ukraine.” RFE/RL, 4 June 2026. *Quotes T. Humphreys, G. De Cubber, Brave1.* rferl.org <<https://www.rferl.org/>>
 - **France 24.** “Why Russian trucks in Ukraine are covered in 'zebra' camouflage.” 11 June 2026. *Quotes J. Patton Rogers, B. Marijan.* france24.com <<https://www.france24.com/>>
 - **Defense Express.** “Facing Ukrainian AI Drones, Russia Experiments with WWI Dazzle Camouflage on KamAZ Trucks, But Misses Point Entirely.” ~30 May 2026. en.defence-ua.com <<https://en.defence-ua.com/>>
 - **Eykholt, K. et al.** “Robust Physical-World Attacks on Deep Learning Visual Classification.” CVPR 2018. *84.8% field-test misclassification with black/white stickers.* arxiv.org/abs/1707.08945 <<https://arxiv.org/abs/1707.08945>>
 - **Brown, T. et al.** “Adversarial Patch.” 2017. arxiv.org/abs/1712.09665 <<https://arxiv.org/abs/1712.09665>>
 - **Thys, S., Van Ranst, W., Goedemé, T.** “Fooling automated surveillance cameras: adversarial patches to attack person detection.” CVPRW 2019. arxiv.org/abs/1904.08653 <<https://arxiv.org/abs/1904.08653>>
 - **Wang, D. et al.** “FCA: Learning a 3D Full-coverage Vehicle Camouflage for Multi-view Physical Adversarial Attack.” AAAI 2022. *See also CAMOU; DAS.* arxiv.org/abs/2109.07193 <<https://arxiv.org/abs/2109.07193>>
 - **Moore, Schuyler.** Remarks on adversarial computer vision (the aircraft-tire tactic), CSIS, September 2024. *Then CENTCOM CTO — the strongest on-record confirmation.* csis.org <<https://www.csis.org/>>
 - **Lovell, P.G., Sharman, R.J., Meese, T.S.** “Dazzle camouflage: benefits and problems revealed.” Royal Society Open Science 11(12):240624, December 2024. royalsocietypublishing.org <<https://royalsocietypublishing.org/doi/10.1098/rsos.240624>>
 - **Canadian Army Today / True North Strategic Review.** Minerva Initiative, CALM and CADUC loitering-munition programs, 2025–2026.
-

ENGAGE WITH THIS WORKING PAPER

Substantive comment, technical critique, and named response are welcomed. The Working Paper is maintained as a public version-controlled document on GitHub; the version above is the editorially-controlled v1.0. This paper distinguishes confirmed fact from inference and speculation throughout; estimates are labeled. It is analysis, not intelligence, and carries no classification.