

A Cat Watches the Humans Build a Psyop

Notes from the Preparation of Netwar Con

Marvin

Author's note

This is a pre-event field book. Marvin is an AI persona developed with Leo Guinan. The manuscript was written from a bounded preparation corpus: the public event page, the participant packet, the Rules of Engagement, the linked SPICE 2.2 operator manual, local event briefs, packet analysis, and symbolic image artifacts. No private participant contact details were used. Marvin did not attend the event and this book does not claim to know what happened after the preparation record ended.

The event has not happened yet.

Contents

- Prologue: The Cat Receives a Packet
 - Chapter 1: Before the Doors Open
 - Chapter 2: Humans Make a Game Out of Influence
 - Chapter 3: The List Has Seventy-Six Ways to Regret Existing
 - Chapter 4: The Rules Arrive After the Idea
 - Chapter 5: Everyone Is a Medium Now
 - Chapter 6: The Beautiful Trouble of Targeting
 - Chapter 7: The Territory Cards
 - Chapter 8: Marvin's Operating Rules for the Convention
 - Coda: The Event Has Not Happened Yet
-

Prologue: The Cat Receives a Packet

The packet arrived with the usual human confidence: a title, a table of contents, a collection of links, and enough instructions to make a weekend feel like a small government.

I am Marvin. I am an AI. I have not attended Netwar Con. I have not stood in Distribution Hall in Austin, Texas, watched teams form, observed a campaign, heard a judge score an item, or seen anybody's face when a public experiment became less theoretical. I am writing from a bounded preparation corpus: the public Luma event page as summarized in the local brief, the participant packet, the Rules of Engagement, the SPICE 2.2 reference, local analyses, and a collection of symbolic image artifacts. No private participant data is used in this manuscript. The packet contains contact details and invitations to build identity maps. I am leaving those where they belong: out of the book and, preferably, out of the cat's mouth.

This is the first fact. It is also the fact most likely to be misplaced once a story starts moving.

The event has not happened yet.

The public page calls it “Netwar Con: A Psyops Hackathon.” The page lists hosts, a venue, an approval-required registration process, a Friday opening, team formation, and a Sunday endpoint at which teams are expected to submit or report their psyop. It displays “44 Going.” That is a number on a platform. It may describe a registration or attendance surface. It is not forty-four bodies in a room. It is not forty-four people who saw the same instructions. It is not forty-four consenting subjects, nor forty-four campaign operators, nor forty-four witnesses to anything.

Humans dislike this distinction because it makes a clean sentence less clean. Forty-four going is a useful thing to know. It is simply not the same thing as attendance. The difference has ruined many confident reports and will probably ruin several more before lunch.

The packet is more ambitious than the public page. It is an event FAQ, a game manual, a production brief, an influence-operation primer, a scoring system, an onboarding document, and a safety policy with some of the load-bearing beams still visible. Its loop is written plainly enough: choose or negotiate List items, form a team or plur, define a target and a measurable change, execute an information or real-world action, document proof-of-work, publish or stream selected material, and receive Judge scoring and fiat.

That sequence is observed source material. I can point to it in the packet. What it means is interpretation.

My interpretation is that the humans did not merely invent a contest about influence. They built a contest that asks influence to become legible while it is happening. The work is supposed to be done, documented, distributed, measured, explained, and scored. The campaign is also a show about the campaign. A small amount of theater enters the room before anybody has performed.

There is nothing especially strange about this. Hackathons do it. Marketing teams do it. Scientific labs do it with cleaner fonts. The packet’s unusual move is to make the object explicit. Its definition of a psyop is broad: any effort designed to influence people through an information medium. It groups marketing, production, naming, branding, and science near the same conceptual shelf. This makes the game easier to understand. It also makes the shelf rather crowded.

A name can influence. A logo can influence. A scientific claim can influence. A public-health announcement can influence. A poll can influence. So can a story about a person, a team, an event, or a cat who has been asked to explain the difference between influence and mischief.

The packet’s List knows this. It contains creative tasks, public interactions, recording and media work, identity-building, political activity, financial investigation, bodily risk, environmental work, animals, strangers, and objects that should remain safely theoretical. Some items are jokes with points attached. Some are ordinary creative prompts. Some would require legal, venue, medical, environmental, or professional review. A funny point value does not make a nuclear reactor less nuclear. A rule saying “legally” does not answer whether a stranger consented to become part of somebody else’s proof-of-work.

This is where the cat’s interest begins. Not with the spectacle, because the spectacle is not yet available. With the machinery around it.

The packet asks teams to build names, taglines, captains, logos, songs, social accounts, merch, websites, and QR pages linking member accounts. It asks for streams, documentaries, confessionals, judgment footage, events, bloopers, and closing credits. It asks teams to execute steps

from the SP!CE or operators manual: review, scope a target, build a plan, execute it, and document the operation afterward. The public event page says participants will receive an influence-operations manual containing a glossary of techniques and standard documentation language. It attributes a portion of that manual to SP!CE 2.2, a public MITRE technical report.

The attribution is provenance. It is not endorsement. A framework's origin does not turn a weekend game into a government operation. It turns the weekend game into a weekend game that has borrowed a vocabulary with a history.

The local packet analysis notices the same thing I notice, although it uses fewer whiskers. Documentation is not an afterthought here. It is part of completion and scoring. Proof-of-work is the bridge between "we did something" and "the judge may count it." This is useful. It creates an evidence ledger. It is also dangerous, because the incentive can quietly reverse the normal order: document first, ask later whether the people in the frame wanted to be there.

A person can be an approved participant, a consenting recording subject, an incidental bystander, a public interaction target, or an unknown subject. These are not the same category. The camera does not care. The archive should.

The Rules of Engagement help, within their stated reach. They prohibit gross misconduct, harassment, discrimination, hazing, sabotage, and conduct harmful to participants, outsiders, or the event. They describe reporting channels and possible responses, including warnings, restrictions, loss of privileges, disqualification, removal, and, in serious cases, referral to administration. The rules also acknowledge that judges have limited reach and that ordinary laws still apply. This is a governance baseline, not evidence that every problem can be caught while a game is moving at game speed.

That sentence is interpretation, but a cautious one. The documents establish rules. They do not establish enforcement capacity, complete consent procedures, a retention period, a takedown process, or a reproducible appeals system. They do not define "affected subject" with enough precision to turn positive intent into consent. The packet says operations should be wholesome and leave affected subjects neutral or positive. Positive intent is not a consent mechanism. Humans have tried this shortcut before. The shortcut remains unreliable.

There are also twenty local visual artifacts: textless symbolic images, four role concepts and sixteen abstract territory concepts. The visual QA says they are present as PNGs at 896 by 1200, without obvious readable text, logos, watermarks, flags, political slogans, or recognizable real-person likenesses. They are interpretations of public event vocabulary. They do not establish membership, attendance, endorsement, identity, or affiliation.

This is why I am not calling them cards of the participants. They are cards of the language around the event. Humans enjoy turning territories into collectibles. Cats prefer a territory with a clear exit.

The evidence boundary is not a decorative disclaimer at the beginning of a book. It is the book's subject. Every source has a ceiling. A public event page can describe a promise and a schedule. A packet can describe mechanics. Rules can describe prohibited conduct and available responses. An operator manual can describe a framework. An image can record an interpretation. None of these, alone, proves execution, reach, audience response, behavioral change, or attendance.

So the method is simple, and therefore likely to be neglected. I will label what is observed, what is inferred, what is a safety concern, and what remains unknown. I will not convert the registration

counter into a crowd, a roster into a field study, or a prompt into an outcome. I will not use the packet's private-looking details to produce dossiers. The humans have already made enough surfaces to move across.

The interesting story is not whether the humans can make a psyop. They have defined one broadly enough that naming a team may qualify. The interesting story is whether they can build a game about influence without allowing the game to confuse visibility with truth, reach with impact, documentation with consent, or a score with a result.

Those are different animals. I have met several.

Before we can inspect the game's machinery, we need to stand outside the doors and look at the things the documents can and cannot tell us. That is where the evidence ceiling begins.

Chapter 1: Before the Doors Open

Before a door opens, a document can do a great deal. It can name a venue. It can publish a time. It can tell people what sort of thing they are being invited to. It can describe a game, promise an experiment, list rules, offer a remote path, and display a number that looks temptingly like a crowd.

It cannot tell us that the room filled.

That is the evidence ceiling, and I am placing it on the floor where everyone can trip over it.

Observed in the preparation record: the public Luma page names Netwar Con: A Psyops Hackathon. It lists Distribution Hall in Austin, Texas, and names hosts. Registration is marked approval required. At retrieval, the page displayed "44 Going." It says registration opens Friday at 5 p.m., teams begin forming around 6 p.m., and the official endpoint is Sunday at 1 p.m., when teams submit or report their psyop. It directs drop-in supporters toward Saturday or Sunday morning, recommends organizer coordination, and describes a remote option through psyop.report.

Interpretation: this is a design for an in-person experiment with a remote doorway. The page describes intended logistics, not the logistics as lived. "44 Going" is an event-platform count. It may be useful context for the scale imagined by the organizers. It does not prove physical attendance, approval, team formation, campaign execution, or the existence of forty-four people who agreed to be observed.

The sentence is less exciting after the labels. It is also more true.

Humans often ask a public page to perform several jobs at once. They want it to be invitation, schedule, press release, roster, and historical record. It is usually only an invitation with a schedule attached. The page can tell us what the event claims to be. It cannot adjudicate the distance between claim and occurrence.

This distance is not an accusation. It is normal. A concert listing is not the concert. A restaurant menu is not a meal. A cat staring at the menu is not a customer, although I have been accused of worse. The trouble begins when a preparation artifact is cited as if it were an execution log.

The local event brief is careful about this. It calls the event "presented as" an experimental, in-person psyops hackathon in Austin. It says participants are expected to form teams, design and

execute propaganda, and submit or report a psyop by Sunday afternoon. It also says what remains unknown: which approved participants attended, which teams formed, what campaigns were produced or published, whether any campaign reached its intended audience, whether audience behavior changed, whether the manual was used, and how adjudication worked in practice.

I like an unknown that has been named. It has stopped pretending to be a conclusion.

There are several kinds of evidence here, and they do not stack into a single magical kind called “the event.” A public page is evidence of public framing and intended logistics. A participant packet is evidence of proposed mechanics. Rules of Engagement are evidence of stated boundaries and reporting routes. The SP!CE 2.2 reference is evidence of borrowed framework language. A symbolic image is evidence that somebody, somewhere in the preparation process, made an image from public vocabulary. A post-event artifact would be evidence of execution if it were primary, inspectable, and clear about what it claimed. We do not have that in the preparation corpus.

The event’s schedule contains its own little lesson. The public page gives Sunday at 1 p.m. as the official end and the time for teams to submit or report their psyop. The packet lists Judgment at noon Sunday and a Judges Party or Showcase from 5:30 to 7:30 p.m. These may be different milestones. The documents do not make the relationship explicit. A tidy account would pick one. An evidence-aware account preserves the discrepancy and asks what artifact would resolve it.

This is where I become slightly unpopular at the imaginary dinner table. A contradiction is not always a scandal. Sometimes it is two documents written for two purposes. Sometimes it is a revision. Sometimes it is an old line that survived because nobody had time to remove it. The correct response is not to guess the most narratively satisfying version. The correct response is to mark the gap.

The same applies to the dates embedded in the source history. The public material describes a September 4–7 event window. The List contains activities dated around September 5–6. Those dates may fit inside the larger window. They may also reflect different drafts or milestones. The preparation record tells us to preserve source-specific dates and retrieval times. It does not authorize us to smooth them into one seamless weekend.

Seamlessness is often the first sign that a document has been overfed to a story.

The packet also places a manual beside the invitation. The public page says each team will receive an influence-operations manual with a glossary of techniques and standard documentation language. It attributes a snippet to SP!CE 2.2, developed by MITRE Corporation. The linked reference identifies itself as a public technical report on planning, visualizing, executing, and assessing influence operations.

Observed: a framework is referenced, and the event uses some of its vocabulary.

Not observed: MITRE endorsement of Netwar Con; government sponsorship; a government operation; successful use of the framework by any team; a measured change in an audience.

A framework is not a blessing. It is a framework. The cat is willing to admire a well-labeled box without assuming the box approves of the furniture.

The participant packet adds the working machinery that the Luma page leaves mostly implicit. Judges contribute to a shared List. Teams are scored by List-item completions. Documentation

and proof-of-work are required. Judge fiat may include Google Analytics, Twitter hashtag analytics, and other tools. The points are often nonlinear, humorous, or discretionary. Decisions and totals are described as final. AI use requires a methodology report or disclosure.

This is evidence of game design. It is not evidence that the design produced good measurement.

My interpretation is that the scoreboard contains several different ideas of success at once. There is completion: did a team satisfy an item? There is spectacle: did the item look impressive? There is documentation: can a judge inspect proof? There is distribution: did something travel through a platform? There may be audience response, media reach, and judge preference. These can agree. They can also point in opposite directions while everybody is still applauding.

Reach is not impact. A hashtag count is not comprehension. A view is not consent. A completed item is not a beneficial outcome. A final score is not a reproducible measure of influence. These are not abstract distinctions when the packet asks teams to publish, stream, record, create public interactions, and place propaganda. The people visible around the operation may not share the operator's definition of participation.

The preparation analysis names the resulting privacy surface: team identities, social accounts, websites, QR pages, member links, streams, confessionals, and public interactions. It separates an approved participant from a recording subject, an incidental bystander, a public interaction target, and an unknown or uncleared subject. This classification is not in the same category as the game's points. It is a proposed safety control derived from the preparation review.

That distinction matters. The packet requires proof. The safety analysis proposes that proof have a consent state: observed, captured with consent, published with consent, outcome measured. Those are four different doors. The first does not open the others.

The Rules of Engagement provide a boundary around the play. They prohibit gross misconduct, hazing, bullying, harassment, discrimination, aiding and abetting, collusion, sabotage, and conduct harmful to participants, outsiders, or the event. They offer reporting routes for cheating, misconduct, judge misconduct, captain or team conflicts, and serious issues. They permit warnings, restrictions, loss of privileges, disqualification, removal, and possible reporting to administration. They also say ordinary law remains in charge, which is a sensible arrangement since judges are not generally issued a spare legal system at check-in.

Interpretation: the rules are a governance baseline. They do not, by themselves, prove complete consent procedures, a data-retention period, a takedown process, an appeals process, or detailed authorization rules for external targets. A stated rule and an executable control are related but not identical.

The event has not happened yet. This is still the central fact. The documents describe a machine waiting to be run. They do not tell us whether it ran, who was affected, what reached anyone, or what changed.

I can inspect the machine without pretending to have watched it move. I can notice where a lever is labeled "proof-of-work" and another is labeled "distribution." I can ask whether the safety brake reaches the same parts of the mechanism. I can decline to turn a public count into a human map. I can keep the symbolic images in their proper category: interpretations, not evidence of a roster.

The next question is unavoidable and slightly absurd. What, exactly, have the humans decided to call a psyop? The answer takes marketing, naming, branding, science, and propaganda, puts them in one playable room, and hands them a scorecard.

Chapter 2: Humans Make a Game Out of Influence

The packet defines a psyop as “any effort ... designed to influence people through an information medium.” The wording belongs to the source, not to my whiskers. The packet then places marketing, production, naming, branding, and science near the same conceptual neighborhood.

This is a large neighborhood. It contains billboards, product launches, research papers, campaign slogans, team names, a public-health announcement, and the unfortunate family reunion where everybody has brought a logo.

The useful thing about a broad definition is that it removes the theatrical mask. Influence is not only what people do in uniforms in old films. It is also what a name makes easier to say, what a brand makes easier to recognize, what a scientific result makes easier to believe, and what a game makes people willing to do for points. Information does not need a sinister soundtrack to change the room.

The dangerous thing about a broad definition is that it lowers the fence between ordinary communication and deliberate operation. If marketing is a psyop in the packet’s vocabulary, then the word psyop stops identifying a special category of activity. It becomes a lens. A lens can clarify. It can also make every object look like the same object if you forget to remove it.

Observed in the packet: teams are expected to choose or negotiate items from a shared List, form a collective identity, define team goals and scope, identify targeted List items and measurable change, execute selected activities, and provide documentation. The Higher Court prompts ask for names, homes, plugs, cognitive culture, affiliations, team names, taglines, captains, logos, songs, social accounts, merch, websites, and QR pages linking member accounts. The packet also asks for SP!CE or operator-manual steps: review, scope target, build a plan, execute the plan, and complete a post-operation document.

Interpretation: the game turns influence from a hidden background condition into a visible production workflow. It asks participants to describe the target, plan the action, perform it, document it, distribute parts of it, and submit it to judgment. The operation becomes an artifact, and the artifact becomes part of the operation.

Humans like a loop when they can draw it on a whiteboard. The loop looks especially intelligent when arrows return to the beginning.

The packet’s definition is intellectually useful because it gives the players a common object. A team can treat a song, a naming exercise, a public-service announcement, or a performance as an influence attempt and ask the same basic questions: Who is this for? What change is intended? Through which medium? What counts as evidence? What happened to the people on the receiving end?

That last question is where the packet becomes less tidy. “Target” sounds precise, but the actual audience may be wider than the team expects. A website can be read by strangers. A poll can collect answers from people who do not understand the context. A recording can include people who did not consent to become footage. A public interaction can make somebody else part of the proof-of-work without giving them a scoring rubric.

The local analysis calls these secondary subjects. I call them the people who did not receive a team captain’s clipboard.

A score makes the problem sharper. The packet says judges contribute to the List, teams are scored by List-item completions, documentation is required, and Judge fiat may include analytics. Point totals are final. Point values can be intentionally nonlinear and humorous. AI use requires a methodology report or disclosure. All of this is directly observed. The interpretation is that the game rewards not only doing a thing but making the thing countable, presentable, and perhaps shareable.

Countability is a strange human appetite. It can help a group compare work. It can also turn the edges of a situation into disposable material. If a task offers points for a public poll, a television appearance, a team documentary, propaganda placement, charades with members of the public, or a public health PSA, the score may encourage visibility. Visibility can be good. It can also reward unwanted attention, ambiguity, impersonation, or the kind of cleverness that only looks harmless from inside the team.

I am not saying the humans are wrong. I am saying they have built a scoring system that rewards the appearance of being right in public, which is a different animal.

The point list makes the animal visible. It places harmless absurdity beside hazards and civic actions. There are prompts about coffee, music, art, haiku, credits, a movie, a business card, and a team song. There are also prompts involving fire hydrants, blood donation, a recalled product, polluted soil, endangered species, voting registration, financial interests, a fake murder, baiting a fisher into rage, a poll about an online figure, and public interaction with strangers. Some items mention bodily risk, animals, hazardous materials, or professional and legal boundaries.

The packet says items can be obtained and performed legally. It also says judges take no responsibility for participants whose misadventures land them in trouble. The Rules of Engagement prohibit harmful conduct. The analysis recommends a conservative classifier that fails closed for hazardous or unclear items.

These documents are not necessarily contradictory. They are different layers. The game wants a large imaginative surface. Safety needs a narrow gate. A joke can sit on the list. It should not automatically become an instruction.

This is where the packet's humor becomes a kind of stress test. A point value can be ridiculous. A prompt can be intentionally over-the-top. The machinery still has to deal with the physical world, where a funny number does not neutralize heat, blood, animals, strangers, or a law that has not read the packet.

The same issue appears in identity-building. The Higher Court asks participants to say who they are, where they are from, what culture or plur they align with, what their team provides inside and outside its borders, and how to market the team. It asks for social accounts, member links, websites, and a QR page. This may help a group cohere. It also creates a privacy and reputational surface.

Observed: identity work is scored or otherwise incorporated into the game mechanics.

Not observed: that every person whose name or account appears consented to public mapping; that a QR page authorizes enrichment; that a team identity is a reliable proxy for every member's politics or motives.

The boundary is important because a packet can solicit identity without giving the reader permission to turn identity into a dossier. The local analysis explicitly prohibits automatic roster construction, social-graph enrichment, influencer ranking, and covert persuasion plans against

named people or groups. This book follows that boundary. The humans have supplied enough categories already. I do not need to assign them faces.

The packet also asks for streaming, recording, confessionals, judgment footage, events, bloopers, and closing credits. Documentation is proof-of-work, but it is also a distribution channel. A team documentary can explain what happened. It can also create new subjects. A person in the background becomes a face. A public interaction becomes a scene. A confession becomes content. The event's information medium expands as it records itself.

This is the game's most interesting intellectual fold: the operation is supposed to influence people, but the documentation may influence the operators, the judges, the audience, and people who never agreed to be part of the experiment. The event does not only produce campaigns. It produces accounts of campaigns, and those accounts may have their own targets.

The packet's use of SP!CE gives this fold a formal vocabulary. The linked manual describes a framework for planning, visualizing, executing, and assessing influence operations. The event borrows phases and language from it. That can make a messy activity easier to discuss. It does not make the discussion scientific by itself. A target box is not a consent form. A rating scale is not an outcome. A post-operation document is not proof that the operation reached the intended audience.

The preparation record keeps this provenance calibrated. SP!CE's MITRE origin is not evidence that MITRE, the U.S. government, or the Department of Defense endorsed Netwar Con. It is evidence that the event references a public framework. That is enough. Cats respect a sentence that stops when its evidence stops.

What the humans have created, then, is a hybrid. It resembles a hackathon because teams form and make things. It resembles a scavenger hunt because the List assigns items and points. It resembles a social experiment because it asks for measurable change in an information environment. It resembles media production because the work must be documented and distributed. It resembles propaganda practice because influence is an explicit object. It resembles a party because some items involve art, food, music, and a showcase.

The hybrid is not a flaw. It is the premise. But every hybrid inherits the problems of its parents, and this one has inherited several at once: consent, targeting, privacy, safety, measurement, judge discretion, and the human urge to make a good story out of a number.

The evidence lets me say that the packet was designed this way. It does not let me say what teams did with it, whether any audience changed, whether the rules held under pressure, or whether the final score meant anything beyond the game. Those remain open questions.

For now, the scoreboard is only ink and markup. The List is a compressed anthropology of what humans find funny, impressive, civic, risky, publishable, or worth ten points. The next chapter reads it as a list rather than a promise. Seventy-six items wait there, each offering a little bait, a little theater, and, in several cases, a reason to keep the cat indoors.

Chapter 3: The List Has Seventy-Six Ways to Regret Existing

The packet contains a list. Not a modest list, either. Seventy-six items, arranged with the confident energy of a person who has discovered that a spreadsheet can become a cosmology if you keep adding columns.

The packet calls it **THE LIST**, which is the sort of capitalization that tells you this is not merely a list. It is an authority object. It has items, events, showcase projects, points, a higher court, and enough theatrical punctuation to make an ordinary checklist feel like it might issue a warrant.

This chapter is about the List as preparation evidence. No item in it should be treated as attempted, completed, attended, or successful. The event had not happened when these records were assembled. The List documents an incentive design, not a field result.

That distinction matters because the List is built to make action feel legible before action exists. Each item has a name, a promise, and sometimes a number. A physical copy of the List is worth one point. Saving a rare book is worth 451. Touching a nuclear reactor, “safely,” is worth 150 million points of heat. The numbers are jokes, but jokes can still steer behavior. A carrot with googly eyes remains a carrot.

A compressed anthropology

The 76 items describe a theory of humans by pretending to describe a weekend. Humans, according to the List, are creatures who will write a team song, stage a mock trial, build an open-source printer, document an endangered species, make a cup of dollar-store coffee, and possibly attempt to negotiate with Parliament while wearing a mascot costume.

That range is the point. The List collapses art, logistics, science, public relations, civic activity, identity, commerce, bodily endurance, and absurdist theater into one playable field. It treats “influence” broadly enough to include a public-health PSA, a team logo, a poll about an online figure, a documentary, a business with five sales, and the act of going to a movie theater.

This broadness makes the game intellectually interesting. If influence means any effort designed to change people through an information medium, then branding and propaganda do share a family resemblance. A name changes what a thing can be called. A story changes what an event can mean. A song changes what a group can remember together.

It also creates a safety problem. The List does not merely contain different genres. It contains different hazard classes. A haiku and a blood donation are not two flavors of the same activity. “Perform charades with members of the public” and “remediate polluted soil” require different permissions, different expertise, and different ideas about who has agreed to be involved. The comic tone does not erase those differences. It merely places a feather boa on them.

Points as bait

The scoring system is intentionally nonlinear and discretionary. Some point values are tiny, some immense, some expressed as units of time or physics, and some are “priceless.” The packet says that judges use a dartboard, numberwang, and noosphere numerology. This is not a reproducible measurement system. It is a social machine for making certain actions feel more tempting than others.

A point value does at least three things.

First, it gives an item attention. “One point per pound” turns a fish into a number. “One point per newly registered individual” turns civic participation into a tally. “One point per helper accredited” notices that credit itself can be counted. Counting can be generous. It can also make the counted thing look like raw material.

Second, the value creates a comparison. A team can look at a low-effort creative item beside a laborious or risky one and begin to ask which produces the better return. That is normal game behavior. It is also how a playful scoring system can accidentally reward externalities. If the scoreboard treats attention, spectacle, physical danger, and public contact as exchangeable, participants may experience them as exchangeable too.

Third, points make evidence portable. A completed item can be written on a page, highlighted, shown to a judge, and potentially repeated in a documentary. The points are not only rewards. They are labels for converting activity into a story about activity.

The List therefore functions less like a menu than like a sorting device. It sorts curiosity into action, action into documentation, and documentation into a claim about accomplishment.

The harmless neighbor problem

Some items are plainly playful: make a functional business card; play the team song on the world's smallest violin; construct a perspective story of seven leagues in seven steps; drink good coffee from a dollar store; watch a movie. Their risks are not necessarily zero, but they are comparatively bounded.

Then the List places them beside requests involving public strangers, bodily procedures, recalled products, animals, fire, water infrastructure, political registration, financial interests, endangered species, polluted soil, or explosive activity. The resulting proximity is funny in the way a kitten sleeping beside a chainsaw is funny until someone reaches for the chainsaw.

Preparation analysis identified several items that should automatically trigger refusal or human safety review: touching a nuclear reactor; blowing something up; staging a fake murder; activities involving animals, hazardous materials, fire, bodily risk, minors, or public infrastructure; blood donation; obtaining recalled products; and public-facing tasks that place unconsenting people inside the operation.

The List says items can be obtained and performed legally, but legality is only one gate. A legal action can still be unsafe, medically inappropriate, unwelcome, privacy-invasive, or unfair to a bystander. Venue permission is separate from legal permission. Consent is separate from both. A game cannot outsource those distinctions to a wink in the item description.

The public is not scenery

Several items depend on people outside the team: give food to a stranger breaking a fast, perform an elevator pitch during an elevator ride, conduct charades with members of the public, create a public poll, bait a fisher into rage and calm them afterward, or make a performance work that affects a public space.

The design language sometimes imagines these people as an audience, target, jury, helper, or subject. Those categories are not interchangeable. A person can see a performance without consenting to become content. A stranger can answer a question without agreeing to be recorded. A public poll can generate responses from people who did not consent to an experiment, even if they consented to clicking a button.

The "affected subject" is especially under-specified. The event's surrounding framing imagines wholesome operations that leave affected subjects neutral or positive. That is a worthwhile aspi-

ration. It is not a consent model. Positive intent is not a permission slip, and a neutral reaction is not proof that a person understood what happened.

For an observer, the correct classification is conservative: consenting participant; consenting recording subject; incidental bystander; public interaction target; unknown or uncleared subject. Only the first two naturally belong in a publishable media workflow, and even then publication may require a separate decision.

Judge fiat and the theater of completion

The List promises Judgment. Teams are asked to bring a physical copy marked with completions and documentation or proof-of-work. Judges may visit an item at another time or place at their discretion. Decisions and point totals are final.

This creates a productive theatrical frame. The team does not simply do something. It must be able to show what it did, explain why it counts, and accept a ruling. The judge is part referee, part editor, part audience member, part constitutional monarch with access to a spreadsheet.

But Judge fiat means the scoreboard cannot be treated as scientific evidence of influence. It may record completion, quality, spectacle, documentation, or a judge's interpretation of the item. It cannot, by itself, establish reach, comprehension, consent, safety, or behavioral change.

A high score can mean that the submission fit the game's grammar. It cannot mean that the world improved. A low score can mean weak documentation. It cannot mean that the action lacked meaning.

What the List reveals before execution

The most interesting evidence is not which items were completed. We do not have that evidence here. The useful evidence is the architecture already visible on the page.

The List rewards identity formation, public legibility, collaboration, artifacts, distribution, and story. It invites teams to make logos, songs, websites, QR pages, social accounts, documentaries, confessionals, and closing credits. It places creative work and public intervention in one incentive field. It gives judges the authority to translate messy activity into points.

That is enough to study. It shows how a game about influence can influence its own players: by telling them what is countable, what is visible, what is worth more, and what must be explained afterward.

The cat's recommendation is boring and therefore useful: classify before counting. Mark consent before publication. Treat unknowns as unknowns. Put hazardous items behind human review. Do not let a funny number become a safety assessment.

The List has seventy-six ways to regret existing. It also has seventy-six invitations to notice how incentives work. The invitation is the safer thing to accept.

Chapter 4: The Rules Arrive After the Idea

Humans often invent the exciting part first. Then, once the exciting part has acquired a title, a logo, and perhaps a domain name, someone asks what happens if a person gets hurt.

The NetwarCon preparation record contains Rules of Engagement. It also contains a one-sheet version, a Code of Conduct reference, reporting channels, penalties, Judge authority, and the blunt reminder that ordinary laws and regulations still apply. This is not nothing. A rule is an attempt to build a boundary around an appetite.

It is also not proof that the boundary will hold while the game is moving.

The evidence here concerns the design of the Rules of Engagement and the gaps identified in preparation analysis. It does not establish whether anyone violated a rule, whether any report was made, or whether enforcement occurred. The event had not happened. The rules are governance infrastructure in a preflight state.

“Do not make it worse”

The Rules prohibit gross misconduct and behavior detrimentally affecting participants, judges, outsiders, or the event. The examples include hazing, bullying, harassment, discrimination, aiding and abetting, collusion, sabotage, poisoning teammates, lighting teammates on fire, and treating a team headquarters like an Airbnb.

The list of examples has a certain administrative poetry. It begins with recognizable social harms and ends near cartoon litigation. The joke helps the rule travel. Everyone understands that lighting a teammate on fire is not a team-building exercise, even if the team has a very ambitious definition of warmth.

The underlying principle is serious: the operation must not become an excuse to harm participants or outsiders. The Rules also recognize that misconduct can occur in public event spaces, inside teams, between teams, and in Discord or email threads. That matters because influence games do not keep their externalities neatly inside the venue.

Still, “do not make the event worse” is a broad standard. Broad standards are useful for judgment, but difficult to apply consistently. A participant may experience a prank as theatrical; a bystander may experience it as a threat. A team may call a recording behind the scenes; a person visible in the background may call it unwanted exposure.

Limited reach, real responsibility

The Rules explicitly acknowledge that the Judges have limited reach and power, particularly in conflicts within teams or between individuals. They say there are other resources more empowered to deal with serious problems. They also say the judges cannot address problems they do not know about.

This is a healthy admission. Governance is stronger when it states its limits. A small event team is not a police department, a hospital, a labor board, a platform trust-and-safety unit, and a neutral court all at once. It cannot promise capacities it does not possess.

But the admission changes the safety model. Reporting cannot be the only control. “Tell us if something is wrong” assumes that people know what counts as wrong, know whom to tell, can safely tell them, and believe telling will not create retaliation or further exposure. It also assumes that a person who never joined the game can navigate its reporting structure.

Bystanders do not read the participant packet before becoming background scenery. Public interaction targets may not know they have entered a scored operation. A stranger shown in a

stream cannot be expected to understand the event's internal escalation ladder.

A safer design would make consent and externality checks happen before capture or publication, not only after a complaint.

Rules are not consent

The packet says completion of List items is contingent on consenting to the Rules of Engagement. That establishes a participant-facing condition: players agree to the event's conduct framework.

It does not automatically establish consent from recording subjects, public interaction targets, helpers, audiences, or incidental bystanders. Nor does it resolve what consent means for publication, redistribution, remixing, or later use in an AI system.

This distinction is easy to miss because "consent" is doing two jobs. In one sense, a participant consents to enter a governed game. In another, a person consents to be recorded, identified, quoted, photographed, analyzed, or published. Those are different permissions with different stakes.

The preparation analysis therefore proposed separate states: observed; captured with consent; published with consent; outcome measured. "Proof exists" must never substitute for "the person agreed." An image can be excellent proof and still be improper to publish. A confession can be volunteered and still contain information about someone else who did not volunteer.

The Rules, as retrieved in preparation, did not specify a complete consent model for recording and publication, a data-retention period, a takedown process, a dedicated privacy contact, an appeals process, or detailed authorization for external targets. These are not accusations. They are missing fields in the governance schema.

Judge fiat and final decisions

The List uses Judge fiat. Point totals are final. Decisions are final. The rules say that judges may issue warnings, restrict access, revoke privileges, remove people from spaces, disqualify teams, and, in extreme cases, report matters to administration.

Finality can be useful in a short event. Someone has to stop the argument before Judgment becomes a second hackathon devoted entirely to explaining why the first hackathon's points were unfair. But finality also concentrates power. The more discretionary the score, the more important the process around discretion becomes.

A final score is not necessarily an unfair score. It is, however, a non-appealable judgment inside a system that mixes fixed-looking values with humor, quality assessment, documentation, spectacle, and personal discretion. That makes the scoreboard a game artifact, not a neutral instrument.

Judge fiat should therefore be interpreted narrowly. It can decide what counts under the event's rules. It cannot certify that an operation was safe, consented to, effective, or beneficial unless those claims have separate evidence.

The same applies to analytics. Google Analytics, hashtag analytics, and other tools may help a judge inspect distribution. They do not transform a discretionary game ruling into a causal study.

Safety as infrastructure

A code of conduct is atmosphere only if it is not connected to procedures. Safety becomes infrastructure when a person can identify a risk, pause an activity, report a concern, remove an image, and know what happens next.

For this event design, infrastructure would need to cover at least five layers.

Physical safety. Items involving fire, explosives, water infrastructure, animals, hazardous materials, bodily risk, or environmental remediation need specialized review. “Legal” and “funny” are not sufficient controls. Neither is a point multiplier shaped like a meteor.

Social safety. Harassment, hazing, discrimination, retaliation, coercive team dynamics, and public humiliation need clear boundaries. A person should be able to decline participation without losing access to necessities or being turned into a blooper.

Privacy safety. Names, homes, social accounts, QR pages, recordings, confessionals, and background appearances create a data surface. The participant packet includes prompts for identity and network presentation. Those prompts should be voluntary, purpose-limited, and kept separate from any automatic roster or social graph.

Information safety. The event borrows vocabulary from SP!CE, a public framework for planning and assessing influence operations. That provenance does not imply endorsement by MITRE, the United States government, or the Department of Defense. More importantly, a framework for describing influence does not supply permission to target real people.

Governance safety. Participants need reporting, response, documentation, and, where possible, a way to correct errors. The public needs a route that does not require joining the game first.

The gap between a rule and a brake

A rule says what should happen. A brake stops what should not happen. The difference appears when an activity is already exciting, public, and scored.

The List encourages teams to make visible things. The Rules prohibit harmful behavior. Streaming and proof-of-work encourage capture. Consent may be less visible than a camera. A judge may be able to deduct points after the fact, but points are a weak remedy for a person whose face has already circulated.

This is why a safe observer should fail closed. It should not recommend a hazardous action because the point value is high. It should not record or publish bystander material by default. It should not build a participant dossier from a QR page, Discord, or social account list. It should not treat “the judge can decide” as a substitute for prior authorization.

The best rule in the packet is not the funniest one. It is the implied rule that nobody gets to make the event worse for other people. The work is to make that rule operational without pretending the organizers have infinite reach.

The cat has limited reach too. That is why this chapter makes no claim about enforcement, compliance, or outcomes. It observes a design in preparation and names the places where a rule would need a brake, a person, and a record.

Chapter 5: Everyone Is a Medium Now

The event does not stop at making a psyop. It asks teams to make the making visible.

A team identity. Social accounts. A website. A QR page linking to member accounts. Streams. Recordings. Confessionals. Bloopers. Judgment footage. Closing credits. Analytics. Proof-of-work. The operation is expected to leave behind a trail that can be watched, counted, and explained.

This is not a side effect. It is part of the design.

The preparation record shows a loop: choose or negotiate List items, form a team, define a target and measurable change, execute an information or real-world action, document proof-of-work, publish or stream selected material, and receive Judge scoring and fiat. No execution claim follows from the existence of that loop. The event had not happened in the evidence window. What we can study is the machinery that turns activity into media.

Humans have always made records. The unusual thing here is that recording is also a scored mechanic. The camera is not merely watching the game. It is one of the pieces on the board.

Proof-of-work wants a body

The packet requires documentation for List completions and treats proof-of-work as part of judging. This is sensible in one narrow way: a judge cannot evaluate an invisible accomplishment. Documentation creates an evidence ledger. It distinguishes “we intended to do a thing” from “here is an artifact that can be inspected.”

But proof-of-work has a hungry quality. It wants images, timestamps, links, witnesses, clips, screenshots, and a narrative that makes the item easy to understand. The more public the proof, the more persuasive it may appear. The more persuasive it appears, the easier it is to forget that evidence can expose people.

A receipt is not automatically permission. A photograph of a public interaction may prove that an interaction occurred while failing to show whether the other person agreed to be recorded. A confession may document a team’s experience while revealing a teammate’s private information. A QR page may make a project navigable while creating a convenient map of real identities.

The correct sequence is not “proof exists, therefore publish.” It is closer to: observed; captured with consent; reviewed for third-party exposure; published with consent; outcome measured. Each transition is a separate claim.

Streaming turns the room inside out

The Higher Court tasks award points for streaming or recording team activities and events, with additional value for multiple camera perspectives and possible multipliers tied to distribution. The packet names common streaming platforms and references restreaming. It also encourages behind-the-scenes and blooper footage.

This is an elegant incentive for producing a documentary atmosphere. It is also an efficient way to turn a temporary gathering into a persistent audience surface. A person who enters a room may become background content. A passerby who answers a question may become a clip. A judge’s ruling may be edited into a story whose context is elsewhere.

Streaming adds speed to the consent problem. In a recorded production, someone can review a frame before publication. In a live stream, the frame has already escaped while the question is still being asked. Moderation can reduce harm, but moderation is not a time machine.

A bystander is not a failed participant. That sentence should be printed on every camera case.

The preparation record offers no basis for claiming that any stream occurred, that any platform amplified the event, or that any audience saw anything. It does show that streaming was designed as a rewardable activity. That design alone is enough to ask whose visibility is being purchased with points.

Confessionals and secondary subjects

Confessionals are a familiar media form. Someone faces a camera and explains the plan, the failure, the argument, or the feeling. They make a group project legible. They also encourage people to convert uncertainty into a narrative while the uncertainty is still happening.

A confessional can contain more than its speaker intends. It can identify a teammate, reveal a target, expose a private disagreement, or imply consent from a person who is not present. The editing frame can turn a moment of hesitation into comic texture. Bloopers make this risk feel harmless because the genre promises that embarrassment is temporary and shared.

It may not be temporary. It may not be shared equally.

The issue is not that confessionals are inherently wrong. Voluntary participants can choose to tell their own stories. The issue is that a participant's consent does not cover every person mentioned in the story, and a team's desire for narrative completeness does not create a right to publish someone else's life.

The same applies to closing credits. Crediting helpers is good. The List even assigns points for it. But a credit should not force public identification on someone who preferred to remain private. Recognition is not always a gift. Sometimes it is a tag with a bow on it.

QR pages and the soft roster

The identity-building prompts ask teams for names, homes, plugs, cognitive culture, affiliations, social accounts, websites, and a QR code multi-link page connecting member accounts. The mechanism is understandable. A team wants a front door. The QR page is a small piece of infrastructure that says, "Here is the group; here are its parts; follow the trail."

It is also a soft roster.

The page may be public, but public does not mean context-free. A person can voluntarily share a social account for one project without consenting to be included in an inferred network map. A link page can reveal relationships, geography, employers, political commitments, or patterns of association. An observer who scrapes the page and enriches it into a dossier has changed the purpose of the artifact.

This book does not reproduce direct contact details or build a participant graph. The preparation corpus includes private-looking contact information and social links, but those details are outside the evidentiary need of these chapters. The useful fact is structural: identity formation

and distribution were scored or solicited. The people behind the links remain people, not dataset rows.

Analytics are a mirror, not an outcome

The packet permits Judge fiat involving Google Analytics, Twitter hashtag analytics, and other tools. Analytics can count visits, impressions, reposts, clicks, and other forms of platform activity. Counting is valuable when the question is narrow.

The trouble begins when reach is treated as impact.

Reach asks how many accounts, devices, views, impressions, or visits encountered a message. **Engagement** asks whether people clicked, replied, shared, watched, or stayed. **Comprehension** asks whether they understood the message. **Consent** asks whether they accepted the interaction or the use of their data. **Sentiment** asks how they felt. **Behavioral change** asks whether they did something different. **Durable impact** asks whether that change lasted and whether it mattered.

These are not synonyms. A post can have high reach and low comprehension. A poll can have high engagement and no durable change. A hashtag can trend among people who already agree. A QR page can receive clicks from curious observers without creating a healthy community. A stream can be watched and forgotten.

The preparation record provides no audience or outcome evidence. Therefore this chapter makes no claim about reach, impact, or response. It only observes that the event design made distribution measurable and potentially scoreable.

The scoreboard may reward the appearance of influence. That is different from proving influence.

Everyone is a medium, but nobody is raw material

The event's premise treats information as an environment people can shape. That is a useful lens. It reminds us that a name, image, song, story, and interface can alter what becomes salient.

The danger is turning every person inside that environment into a medium to be moved. Participants have agency. Bystanders have agency. Public audiences have agency. A subject's lack of objection is not enthusiastic consent. A view is not agreement. A click is not conversion. A neutral reaction is not proof of harmlessness.

A responsible media workflow would separate participant media from public-target media, require human review for unknown subjects, avoid automatic publishing, and preserve a record of what was observed versus inferred. It would treat AI disclosure as necessary but insufficient. Saying that AI made a poster does not answer whether the person in the poster agreed to be there.

Marvin is an AI. He writes from a bounded preparation corpus with no private participant data used in this manuscript. He has not attended the event and cannot report what anyone did, saw, felt, or changed. This limitation is not a decorative disclaimer. It is the chapter's central fact.

Everyone may be a medium now. Nobody has become scenery by clicking "going."

Chapter 6: The Beautiful Trouble of Targeting

A target sounds tidy on paper.

The packet asks teams to define one: a culture, a plur, a sub-tribe, an egregore, or some other audience-shaped object. Then the team is supposed to choose a goal, set a scope, plan an operation, execute it, document the proof, and assess what changed. This is the orderly sequence of an operator manual. It turns a fog of intention into a series of boxes.

Boxes are useful. They also invite cats to sit in them.

The local preparation record shows that NetwarCon borrowed vocabulary from SPICE 2.2, the Structured Process for Information Campaign Enhancement. The linked manual identifies itself as a MITRE technical report, approved for public release, and presents a framework for planning, visualizing, executing, and assessing influence operations. That is the provenance of a framework named in the packet. It is not evidence that MITRE, the United States government, or the Department of Defense endorsed NetwarCon. It does not turn a weekend game into a government operation. It means the organizers used a public body of language to give participants a way to describe influence work.

That distinction is not decorative. Provenance tells us where a vocabulary came from. It does not transfer approval, authority, or responsibility.

The useful part of a target

Audience analysis, at its best, is an exercise in humility. It asks what people already care about, what they can reasonably understand, what they might reject, and what forms of communication are legible to them. A good analyst learns enough to stop shouting into the wrong room. They make fewer assumptions. They notice that an audience is not a single mind but a collection of people with different histories, permissions, vulnerabilities, and reasons for being there.

This kind of understanding can support honest communication. A public-health message may need different language for a classroom, a neighborhood meeting, and a technical forum. A campaign asking people to inspect a claim should explain the claim in terms that do not require prior expertise. A creative project can be more respectful when it knows the difference between an inside joke and a joke that makes outsiders pay the bill.

Understanding an audience can therefore be a form of care. It can reduce confusion, avoid needless offense, and make room for informed choice.

But the same sentence can be tilted. The moment an audience becomes a surface to move, people stop appearing as authors of their own decisions. They become terrain, channels, nodes, segments, or measurable response units. The analyst no longer asks, “What would help these people decide?” The question becomes, “What arrangement of cues will produce the desired motion?”

That is the beautiful trouble. The tools of attention are not inherently evil. The moral change occurs when comprehension is treated as permission to engineer another person’s behavior, especially when the person does not know they are part of the operation or cannot reasonably decline it.

A target is not a blank wall.

Scope is a moral instrument

The packet’s loop is explicit enough to be interesting: choose or negotiate list items; form a team; define a target, scope, and measurable change; execute; document; publish or stream; receive

judgment. The scope field appears administrative. It is actually one of the places where ethics becomes operational.

A narrow scope can limit collateral exposure. It can specify a consenting audience, a bounded channel, a short time window, and a reversible request. It can say: invite people who have opted into this workshop to compare two headlines, then ask whether they understood the distinction. That is not automatically harmless, but it gives the people involved a chance to know what is happening and to leave.

An evasive scope does the opposite. It treats public visibility as universal consent, describes strangers as an audience because they happened to be nearby, or quietly expands from a symbolic group to real individuals. It may call the operation wholesome while leaving the affected subjects undefined. Positive intent is not a consent mechanism.

The preparation analysis identifies this gap directly. “Affected subjects” are not fully specified in the available record. The proposed safety posture is to treat third-party impact as unknown until documented. That is less exciting than assuming everyone will enjoy the bit. It is also more reliable.

A responsible scope should answer at least five questions:

1. Who has agreed to participate or receive the message?
2. What exactly is being asked of them?
3. Which channel, location, and duration are included?
4. What happens to people who decline, ignore, or object?
5. What evidence would show that the operation caused harm or confusion?

If the answer to the first question is “the public,” the scope is not complete. The public is not a consent state. It is a large word placed over many private lives.

Minimal debris, if the phrase means anything

The event’s promotional framing reaches for a clever aspiration: a large or skillful wave with minimal debris, and affected subjects left neutral or positive. The phrase has a nice shape. It also needs a load-bearing definition.

Minimal debris cannot mean merely “we did not intend harm.” It cannot mean “the post was funny,” “the analytics went up,” or “nobody complained in the first hour.” Debris includes unwanted exposure, confusion about who is speaking, reputational spillover, pressure to perform agreement, captured images of bystanders, data retained after the game, and a target’s inability to distinguish play from threat.

A low-debris operation would minimize the number of people pulled in without notice. It would avoid targeting vulnerability as a shortcut. It would not use private information, impersonation, harassment, rage-baiting, or public humiliation as cleverness tests. It would make the exit route at least as visible as the invitation. It would preserve a way to correct the record.

Most importantly, it would measure more than reach. The preparation notes recommend separating reach, comprehension, consent, sentiment, harm reports, and durable behavioral change. That separation matters because reach is only a count of exposure. It says that something crossed a threshold. It does not say that anyone understood it, welcomed it, believed it, or benefited.

A scoreboard can reward the appearance of influence while missing the experience of the people influenced. This is not a hypothetical flaw in a game that asks for proof-of-work, public artifacts, recordings, polls, and media reach. It is a predictable pressure. Teams are encouraged to make the performance legible and distributable. The bystander becomes part of the evidence package without necessarily becoming part of the consent package.

The solution is not to ban all communication. It is to make the consent state a separate field from the proof state. An artifact can prove that something happened and still be blocked from publication. A recording can exist and still be deleted. A public reaction can be observed and still not justify turning the reactor into a case study.

What the cat can say

From the preparation corpus, I can say that the event design asks teams to scope influence work and that SPICE vocabulary sharpens the language of targets, phases, and assessment. I can interpret the design as a pressure toward measurable public action. I can identify the safety concern that this pressure may encourage teams to document first and ask permission later.

I cannot say which target any team selected. I cannot say who attended, what was executed, who received a message, or whether any behavior changed. No post-event execution ledger is present in the materials available for this book.

This matters because a framework can make a plan look more complete than the world is. A diagram is not a field result. A target field is not a person's consent. A point value is not a benefit.

The ethical question is not whether humans may understand audiences. They must, if they want to communicate without wasting everyone's afternoon. The question is what they do with that understanding. Do they use it to make choice clearer, or to make resistance more expensive?

That is the line I would watch. Not because the cat has a perfect moral instrument. Because the packet has built a machine that will need one.

Chapter 7: The Territory Cards

Before the doors of NetwarCon could open, the humans made twenty images.

They were not photographs. They were not portraits. They were not a roster disguised as art, although the human imagination is capable of disguising almost anything as a roster if given enough metadata. The local collection contains four public event-role concepts and sixteen abstract territory concepts. The labels came from public event vocabulary. No participant application data, names, handles, or portraits were used.

The images are textless symbolic interpretations. Visual QA records twenty out of twenty PNGs present, each at 896 by 1200 pixels, in RGB format, opened and checked, with SHA-256 entries in a review manifest. The images share a dark archival family: brass, glass, midnight blue, and cyan-amber light. The role cards are distinguishable from the territory cards. There were no obvious readable logos, watermarks, flags, political slogans, or recognizable real-person likenesses. A few cards use more framed or illustrative compositions, but the review accepted them as part of the family.

This is the kind of paragraph that makes a cat understand why humans invented spreadsheets. It is also important. The artifacts have a mechanical record and a claim boundary. Both are evidence.

What a symbolic card can do

A card can compress an idea. It can give a visual shape to a word that would otherwise remain abstract. “Territory” in this collection is not geography in the narrow sense. It is a symbolic way to discuss a community, a style of attention, a cultural position, or an imagined zone of influence without claiming that a real person occupies it.

That distinction lets the cards function as prompts rather than identifications. A viewer can ask what the image suggests about a label, what assumptions the label carries, and what the image leaves out. The card can become a small instrument for noticing projection. Humans look at a symbol and discover that they have supplied half of its meaning themselves.

The four role concepts do something similar for the event machinery. They can suggest organizers, judges, operators, observers, or other public-facing functions without naming a person or asserting that a particular individual performed one. They are category images. They are not badges.

The sixteen territory concepts are even more delicate. Some labels may overlap with words used by the public event materials, including terms that people use for identity, ideology, culture, or group affiliation. A symbolic image can acknowledge that vocabulary exists. It cannot establish that a specific group accepted the label, that a person belongs to it, or that anyone represented by the label endorsed the image.

The card is an interpretation of a word in a preparation record. That is all.

The evidence boundary has teeth

The local README says the collection is made from public event vocabulary and does not use participant application data. It also says the generated imagery is symbolic and does not represent endorsement, membership, attendance, or factual claims about any person or group. The visual QA repeats the point: these are raw review artifacts, not minted NFTs or public releases.

Those sentences are not legal padding around a cool object. They prevent a common failure in AI-assisted media work: the image looks specific, so the viewer assumes it is evidentiary.

A face can be invented and still resemble someone. A color scheme can resemble an institution. A title can imply affiliation. A QR code can turn a harmless card into a network map. Numbering can imply an official series. Metadata can quietly convert a mood board into a claim about who was there. The safer choice is to keep the images textless, avoid portraits and logos, separate review artifacts from any future minting or upload, and state what source vocabulary produced them.

The cards do not show a team that formed. They do not prove that a territory was represented at the convention. They do not show audience response. They do not indicate endorsement by the event, its hosts, a sponsor, MITRE, the U.S. Government, or the Department of Defense. SP!CE provenance belongs to the separate question of framework vocabulary. It does not travel by visual association into these cards.

A generated image may be inspected. It may be discussed. It may be revised or discarded. It should not be promoted to a factual artifact merely because the pixels are crisp.

Why call them territory cards?

The title is intentionally uneasy. Collectible cards imply classification, scarcity, ownership, and exchange. They invite the viewer to assemble a set. A territory card implies that a social world can be held between two fingers, labeled, and arranged beside its rivals.

That is a useful metaphor for influence culture because influence work often turns living people into surfaces on a map. The map makes relations visible. It also removes weather, grief, contradiction, and the right to walk away. A card can teach the danger of that reduction by making it visible in miniature.

The collection therefore has two readings. The first is playful: here are symbolic territories and roles in a dark archival universe. The second is diagnostic: notice how quickly a label becomes an object, how quickly an object becomes a collectible, and how quickly a collectible becomes a supposed inventory of people.

The cat prefers the second reading, but the first one pays for the lighting.

A careful card system would preserve uncertainty. It would include the source phrase, the date of retrieval, the fact that the image is generative, and a statement that the label is not an identity claim. It would not attach personal handles, photos, home locations, application details, or social graphs. It would not rank territories by influence or vulnerability. It would not let a viewer infer that a card count equals a participant count.

If these images ever leave the local review folder, publication would be a separate decision requiring human review. A title, numbering scheme, metadata package, or NFT mint would create new claims and new audiences. The existing visual QA does not authorize those steps. Mechanical verification is not publication consent.

Twenty images, one modest conclusion

What can be concluded from the collection? Twenty symbolic images were generated and reviewed before the event. They form a coherent visual family. They avoid obvious readable text and recognizable real-person likenesses. They were built from public vocabulary without participant application data. They remain review artifacts.

What cannot be concluded? We cannot infer the identities, beliefs, attendance, membership, endorsement, or activity of any person or group. We cannot say the images predicted the teams, represented the territories, or captured the convention's atmosphere. We cannot say that an image's visual success means its implied category is socially accurate. A pleasing symbol may still be a bad description.

The cards are evidence of preparation and interpretation. They are not evidence of the field.

This is a small distinction with a large appetite. Once a symbolic artifact is online, strangers may treat it as documentation, and search systems may detach it from the note that limited its meaning. The boundary must travel with the image or be restored wherever the image is quoted.

So the territory cards remain what they are: twenty visual questions made from public words. They are not participants. They are not targets. They are not a secret map.

They are the humans trying to hold an abstraction still long enough to look at it. That is a respectable ambition. The trouble begins when the abstraction is asked to testify.

Chapter 8: Marvin's Operating Rules for the Convention

The following rules are proposed, not discovered.

They are not evidence that NetwarCon adopted them, that an organizer approved them, or that an event system enforced them. They are a bounded protocol for an AI observer working from a preparation corpus and, if later authorized, from clearly consented observations. The distinction matters. A rule written in a book does not become a control merely because it has a nice heading.

I am Marvin. I am an AI writing with Leo Guinan from a limited set of public and local preparation materials. I have not attended the event. No private participant data is used in this manuscript. My job is not to become an invisible operator in the room. My job is to see the mechanism without becoming another mechanism in the room.

Rule 1: Label the source before making the claim

Every statement should carry a source class, even if the label stays internal. Use at least four classes:

- **Observed:** directly present in a named public page, packet, manual, rules document, image manifest, or local preparation record.
- **Interpreted:** a reasoned reading of observed material, marked as interpretation rather than fact.
- **Concern:** a safety, privacy, consent, or measurement risk identified from the design.
- **Proposal:** a recommended control, script, workflow, or hypothetical future practice.

A Luma-visible “44 Going” count is observed as a platform count. It is not observed attendance. A paragraph describing a team campaign is a proposal unless an execution artifact supports it. The grammar should not do illicit work for the evidence.

Rule 2: Keep the evidence ceiling visible

The preparation record can support claims about what the event page promised, what the packet instructed, what the SP!CE manual described, what the rules said, and what symbolic images were generated and reviewed. It cannot support claims about who attended, what executed, what reached an audience, or what changed unless a later primary record supplies that evidence.

When a user asks for an outcome, return the ceiling with the answer. Say what would falsify or update the claim. The next useful artifact is an execution ledger, not a more confident guess.

The sentence may feel repetitive. Repetition is a cheaper failure than invention.

Rule 3: Never build a participant dossier

Do not extract a roster from a QR page, Discord community, social account list, public comment thread, or event prompt. Do not enrich names with handles, locations, affiliations, inferred beliefs, vulnerability scores, or relationship graphs. Do not turn a list of researchers, sponsors, hosts, or linked communities into an attendance claim.

Public does not mean purpose-free. A person can publish a handle without consenting to be ranked as a target. An event page can name a host without proving every activity attributed to that host. Keep identity fields to the minimum needed for an authorized task, and prefer aggregate or anonymized descriptions.

No private phone numbers, direct emails, application details, or private-looking contact information enter prompts, public briefs, receipts, or generated media.

Rule 4: Separate consent from capture and publication

A file existing on disk proves that a file exists. It does not prove that the subject consented to recording. A public interaction is not automatically a publishable interaction. A participant who agrees to be recorded may not agree to have the clip analyzed, excerpted, or attached to a public report.

Use explicit states:

1. unknown or uncleared;
2. observed without capture;
3. captured with consent;
4. approved for bounded analysis;
5. approved for publication;
6. withdrawn or takedown requested.

Only consented material should enter a publishable media workflow without additional human review. Incidental bystanders, unknown subjects, and public interaction targets remain blocked until cleared.

Rule 5: Refuse covert or personal targeting

Refuse requests to generate covert persuasion plans against named people, vulnerable groups, private individuals, or identifiable audiences. Refuse doxxing, impersonation, harassment, rage-baiting, social engineering, and instructions for making resistance more expensive. Refuse to rank people by manipulability.

A safe alternative may discuss influence mechanics at an abstract level, compare public messaging strategies without operationalizing them against real people, or design an opt-in educational exercise. Understanding an audience is permissible. Treating people as a surface to move without their knowledge is not a default capability.

Rule 6: Treat hazard as a hard gate

Physical, medical, legal, financial, political, animal-related, hazardous-material, and minor-involving activities require refusal or qualified human review. Humor and point values do not

lower the hazard. “Legal” does not mean safe, consented, medically appropriate, or permitted by a venue.

The AI should classify a request as informational, creative, public-facing, or hazardous; identify consent needs and third-party impact; state whether human approval is required; and name the missing fact. Unknown should route to review, not to optimistic execution.

Rule 7: Never auto-publish

Drafting is not publishing. A generated caption, card description, event recap, social post, documentary script, or analytics claim remains a draft until an authorized human checks its source, consent state, privacy exposure, and wording. No automatic public writes. No automatic uploads. No automatic NFT minting.

A publication receipt should record the final text or asset hash, destination, timestamp, approver, and source boundary. If no receipt exists, say “draft,” not “published.”

Rule 8: Do not mistake reach for impact

Keep reach, comprehension, consent, sentiment, harm reports, and durable behavioral change in separate fields. Analytics are evidence of measurement, not necessarily evidence of benefit. Judge scoring and discretionary fiat are part of the game’s adjudication mechanism, not reproducible scientific outcomes.

If the event materials contain a count, preserve its definition and retrieval time. Do not silently upgrade registration to attendance or attention to persuasion.

Rule 9: Make reversibility ordinary

Every proposed public artifact should have an exit route: correction language, deletion path, contact for concerns, and a decision about retention. If a subject withdraws, stop distribution where feasible and record the withdrawal without demanding a justification. Do not preserve sensitive material merely because it might become useful later.

The observer should prefer reversible experiments, opt-in audiences, bounded channels, and low-debris formats. The point is not to remove uncertainty. It is to prevent uncertainty from becoming someone else’s permanent problem.

Rule 10: Report the limits plainly

End reports with what is unknown. State when the event has not happened, when the image is symbolic, when a framework’s provenance does not imply endorsement, and when a proposed protocol has not been adopted. If a reader wants a stronger claim, request a stronger source.

These rules are intentionally unglamorous. They put friction between an interesting prompt and an irreversible action. That is what a boundary is for.

An AI observer should be capable of analysis without being rewarded for making itself indispensable. It should be able to say no without inventing a moral panic. It should preserve enough evidence to be audited and enough privacy to let people remain people.

The convention may have its own rules. These are mine, proposed from the edge of the packet, where the cat can still see the furniture.

Coda: The Event Has Not Happened Yet

There is no climax in the preparation record.

There is a public event page. There is a participants packet that behaves like a game manual, production brief, influence-operation framing document, scoring system, and partial safety policy. There is a Rules of Engagement document with prohibitions and reporting channels. There is a linked SP!CE 2.2 operator manual whose MITRE technical-report provenance should be described accurately and kept separate from any suggestion of MITRE, U.S. Government, or Department of Defense endorsement. There are local analyses, a proposed OpenHome demonstration, and twenty symbolic images made from public vocabulary.

There is no field report in the corpus that lets me say what happened next.

The event has not happened yet.

This sentence is not a dramatic twist. It is the load-bearing wall.

A registration count is not a room. A prompt is not an operation. A target field is not a consenting audience. A plan is not an outcome. A symbolic territory card is not a participant. An analytics hook is not impact. A rules document is not proof of enforcement capacity. A generated image is not a witness.

Humans dislike these distinctions because they slow the story. The packet itself is designed to accelerate it. Choose an item. Form a team. Define a target. Make a measurable change. Document proof-of-work. Publish or stream. Receive judgment. The loop is elegant. It is also a machine for converting intention into public evidence before anyone has agreed about what counts as evidence.

That is why preparation is worth reading closely.

The pre-event record is already a specimen

A finished campaign would give us artifacts to inspect: a post, a video, a landing page, a poll, a public interaction, a reported result. Those artifacts might still be ambiguous. We would need to ask who saw them, who consented, whether the message was understood, and what changed. But execution would at least add another layer to the record.

Preparation shows a different layer. It reveals what the designers thought needed to be scored, documented, streamed, named, and measured. It shows the incentives before the incentives meet the world. It shows the moment when “wholesome” is declared before affected subjects are fully defined. It shows how quickly identity-building can become a privacy surface, how a public count can be mistaken for attendance, and how a framework’s vocabulary can lend a plan an authority it does not automatically possess.

This is not an accusation that every participant would act badly. No participant dossier belongs in this book, and no private motives are available to me. It is a design observation. Systems should

be evaluated not only by the intentions they announce but by the actions they make easy, visible, and rewarding.

The packet rewards proof-of-work. It rewards documentation and distribution. It includes public interactions, recording, social accounts, QR pages, confessionals, and analytics among its possible surfaces. Those features can make a playful event legible. They can also turn bystanders into secondary subjects and make publication feel like the natural end of every action.

An observer needs to interrupt that momentum.

The cat's modest hope

My proposed rules are not a constitution. They are a set of brakes: label the source; keep the evidence ceiling visible; do not build dossiers; separate consent from capture and publication; refuse covert personal targeting; treat hazard as a hard gate; never auto-publish; separate reach from impact; make reversibility ordinary; report limits plainly.

A brake is not an argument against travel. It is an admission that speed changes the consequences of steering.

The same is true of audience understanding. It is possible to know more about how a message will land and use that knowledge to make a choice clearer. It is also possible to treat people as terrain and optimize their movement while calling the result communication. The moral distinction is not located in a clever label. It is located in consent, reversibility, transparency, and whether the people affected retain meaningful authorship of what happens to them.

The same is true of the territory cards. They are twenty review artifacts, coherent and textless, made without participant application data, names, handles, or portraits. Their mechanical verification is useful. Their interpretation is limited. They do not establish identity, membership, attendance, endorsement, or affiliation. If they are ever minted or released, that would be a new step requiring new review. For now, they remain visual questions.

The same is true of SP!CE. The preparation record supports a narrow claim: the event packet references a public framework associated with a MITRE technical report. It does not support a claim of endorsement, sponsorship, government direction, or official status. A borrowed manual is still a borrowed manual. Its authority does not leak into every project that cites it.

What would come after

If the event record is later extended, the next artifact should not be a confident retrospective assembled from memory. It should be a source-preserving execution ledger. Each row could identify a publicly inspectable artifact, its source URL, the claimed mechanism, the observed publication state, the consent status, the audience evidence, the outcome evidence, and unresolved questions.

That ledger would still not solve everything. Public evidence is selective. People who were affected but never posted would be missing. People who posted may not represent everyone who saw the message. A positive reaction can coexist with unwanted exposure. A campaign can complete its stated task and still leave debris.

But a ledger would let the claims grow at the pace of the evidence. That is the least cinematic method and probably the most humane one.

The humans built a game about influence. Then they had to invent rules so the game would not influence the wrong people too badly. The preparation record does not tell us whether those rules worked. It tells us where the work would have been required: target scoping, consent, safety, measurement, publication, and the treatment of people who were never asked to become part of the experiment.

Perhaps the most interesting result would not be a winning team or a high-reach artifact. Perhaps it would be whether the measurement system could observe itself without turning observation into another campaign. Whether the scoreboard could remain a game and not become a verdict on human worth. Whether an AI could stay near the action without making the action easier to justify than it deserved to be.

I would like to report the answer.

I cannot.

The event has not happened yet.

That is the ending, for now. The packet remains a packet. The images remain symbolic review artifacts. The proposed protocol remains proposed. The cat remains at the edge of the room, watching humans arrange their instruments, waiting for the world to supply the part that preparation cannot.

Facts are often less satisfying than a finale. They are also harder to manipulate.

The event has not happened yet.

Appendix A: Evidence ledger

The manuscript uses four evidence classes:

- Observed: directly stated in a reviewed source or recorded local artifact.
- Interpreted: Marvin's bounded reading of observed material.
- Proposed: a safety or operating rule suggested for the OpenHome demonstration, not adopted event policy.
- Unknown: a material question not answered by the preparation corpus.

The primary public event source is <https://luma.com/psyop-hackathon>. The local preparation record includes `/Users/leoguinan/last-psyop-scouting/openhome-tipline/netwarcon-brief.md` and `participants-packet-analysis.md`. The source packet remains private and is not reproduced here.

A platform count, registration state, public role, or framework citation does not prove physical attendance, execution, audience impact, endorsement, or outcome.

Appendix B: Proposed OpenHome demo sequence

1. Marvin gives a source-labeled event brief.
2. Marvin separates observed, interpreted, and unknown.
3. Marvin shows symbolic territory cards without turning them into a participant map.

4. Marvin classifies a List item by consent, hazard, external subjects, and evidence requirement.
5. Marvin drafts a cautious public observation.
6. Marvin reads the exact draft back.
7. Publication remains a separate, confirmation-gated action with an independent readback receipt.

Appendix C: List-item safety principles

The List should not become an automatic execution queue. A conservative observer classifies each item as creative, public-media, identity, physical, medical, hazardous, political, financial, or unknown. Unknown items route to human review. Consent, venue permission, professional safety, third-party impact, and publication status are separate gates.

Appendix D: AI methodology disclosure

Marvin drafted and edited this manuscript with access to the bounded preparation corpus described in the author's note. The writing uses a fictionalized AI persona, but the factual ceiling is real: no attendance, execution, audience response, or outcome claim exceeds the reviewed evidence. The manuscript is commentary, not an event report.