

NAI GUIDANCE

Key Do's & Don'ts for Using **AI** in Network Advertising

Non-binding guidance for NAI members addressing privacy, transparency, oversight, and accountability issues to enable newer AI-enabled and agentic advertising workflows.

JULY 2026

TABLE OF CONTENTS

A. Introduction..... 3

 1. Purpose & Scope.....3

 2. Relationship to the NAI Framework and other industry work..... 3

B. Key Do’s & Don’ts for AI Use in Network Advertising.....4

 1. Inventory and labeling..... 4

 2. Segment review and activation..... 5

 3. Testing and monitoring..... 6

 4. Disclosure..... 7

 5. Permissions and constraints.....8

 6. Choice and signal handling..... 9

 7. Oversight and logs..... 10

 8. Contracting..... 11

 9. Accountability..... 12

C. Suggested implementation checklist..... 13

D. Notes.....14

A. Introduction

1. Purpose & Scope

This guidance is a voluntary tool for NAI members to help enable the use of AI in Network Advertising by identifying potential privacy and transparency issues to address, especially for audience creation, targeting, optimization, measurement, and emerging agentic advertising workflows. It identifies practical steps members can take now, even while technical standards and market practices continue to evolve. It does not address general-purpose model training, AI-generated ad creative, AI uses outside Network Advertising, or protocol design issues. Not all topics will be relevant to every member, and the allocation of responsibility across vendors and customers will vary depending on each member's role.

In offering this guidance, the NAI recognizes that the use of AI (broadly construed) is not new to Network Advertising. NAI members have been among the earliest and most sophisticated adopters of machine learning, optimization, and predictive modeling for well over a decade, and the [NAI Framework's](#) existing principles already apply to those established practices. This guidance is not intended to add new controls or obligations for established optimization and modeling tools already governed by the NAI Framework.

Instead, it addresses two categories of newer development where additional privacy considerations may come into play:

1. **Generative and broadly adaptive AI systems**, including natural-language audience creation tools and other systems that accept broad or unstructured inputs, infer audiences or attributes beyond human-specified parameters, or materially reduce the ability to review outputs before activation.
2. **Agentic advertising workflows**, where autonomous systems can bid on inventory, access data, and execute transactions across intermediaries without a human approving each step. These workflows introduce new questions about consent chains, data lineage, signal handling, and oversight that did not arise when humans made each decision.

The additional considerations below for newer AI-driven workflows are designed to help members focus on review and oversight resources where they matter most.

2. Relationship to the NAI Framework and other industry work

This guidance helps illuminate how the NAI Framework's principles interact with the use of AI in Network Advertising. It is most directly relevant to the Data Governance principle, while also supporting Transparency, Choice & Consumer Control, Sensitive Personal Data, and Accountability principles.

It is designed to complement, not replace, existing technical privacy and governance resources, and does not endorse any specific protocol, architecture, or standards body. Other industry organizations are producing valuable work on related questions, including legal compliance frameworks for AI adoption, transparency standards for AI-generated ad creative, and privacy architecture for agentic systems.¹

This guidance is provided for informational purposes and does not constitute legal advice. Members should consult their own legal counsel regarding the application of these recommendations to their specific circumstances.

B. Key Do's & Don'ts for AI Use in Network Advertising

The recommendations below apply in proportion to a system's autonomy: the more an AI system can act, particularly where it can change privacy or legal outcomes before a person reviews the specific action, the more carefully these recommendations should be considered.

1. Inventory and labeling

An inventory shows where AI is used across advertising workflows, which controls apply, and who is accountable, before issues arise.

DO

Maintain a current inventory of AI-enabled advertising tools and workflows, and ensure that the role AI plays is appropriately described at each level where it is relevant.

- Track AI-enabled products, features, models, agents, and AI-generated segments used in advertising workflows.
- Record the business purpose, principal inputs and outputs, owner, and deployment status for each system.
- Categorize each system by its level of autonomy and human direction, for example, distinguishing among (a) advisory and analytics tools that surface recommendations for human decision-makers; (b) optimization tools that adjust parameters within human-defined constraints; and (c) generative or agentic systems that construct novel outputs or execute transactions with limited or no human approval at each step. This categorization determines how much of the additional guidance below is relevant to each system.
- Ensure that the role AI plays is described appropriately for each relevant audience. For example, consider whether and how:
 - Internal records should reflect how each system is categorized and which controls apply to it;
 - Advertiser and partner-facing materials (e.g., product documentation, terms, or campaign-set-up interfaces) should describe the functional role AI plays when it meaningfully shapes how data is processed, how audiences are constructed, or how decisions are made. Note that this does not require disclosing model names, architectures, or other proprietary information. Instead, consider how to describe what the AI does in terms a business partner can act on; and
 - Consumer-facing disclosures should, where warranted, reflect material AI use (addressed more fully in Section 4).

DON'T

- ✘ **Don't let AI-enabled systems or outputs operate in advertising operations without considering the inventory and categorization needed to facilitate appropriate review. Systems that were never inventoried or categorized tend to surface only after something has gone wrong, when review options are narrower and remediation may be more difficult.**

2. Segment review and activation

Segment review can help identify potentially sensitive or proxy outcomes before an AI-generated audience is activated.

DO

Review AI-generated segments and their activation context before use, particularly where sensitive topics or populations may be involved.

- Identify AI-generated segments (or AI-enabled segments guided by natural-language instructions) to enable separate review where appropriate.
- Consider reviewing AI-generated segments where the activation context is sensitive (e.g., health, sexuality, precise location, minors, or other sensitive topic areas).
- Consider reviewing the activation outcome. Could it be an inadvertent proxy for a sensitive attribute?
- Maintain a list of sensitive verticals or campaign types (such as health, credit, debt relief, financial distress, housing, employment, or campaigns directed to minors) that may warrant additional review before AI-generated audiences are activated.
- Where natural-language audience creation is enabled, consider constraining the inputs available to the tool. This may include restricting source taxonomies or datasets to exclude sensitive topics, applying input filters that prevent the system from accepting prompts directed at sensitive categories, or both.

DON'T

- ✘ **Don't assume an AI-generated segment is low risk merely because it lacks an explicit sensitive label, or rely solely on whether the segment was intended to reach a sensitive audience. Where natural-language audience creation tools are available, don't leave sensitive data sources or categories accessible to those tools without appropriate guardrails.**

3. Testing and monitoring

Testing surfaces proxy risk, drift, and boundary failures before they can impact consumers, partners, or campaigns.

DO

Test for proxy risk, drift, and other unintended outcomes, at a level proportionate to the system's capabilities and the sensitivity of the context. For higher-risk agentic workflows, where the agent can change legal or privacy outcomes before a person reviews the specific action, consider repeatable test cases or benchmarks to evaluate whether the system stays within approved data-access and retention constraints.

Higher-risk agentic workflows may include workflows where an autonomous or semi-autonomous system can access data, select audiences, route data, bid on inventory, activate campaigns, honor or apply privacy signals, interact with third-party systems, or execute transactions without human approval at each step, particularly where sensitive contexts, cross-party data flows, or external platform access are involved.

- Conduct a reasonable pre-launch assessment to evaluate whether an AI-generated segment could act as a proxy for sensitive categories, even if nobody directed that outcome. A reasonable assessment does not require algorithmic auditing or access to model internals. For lower-risk uses, it may be as simple as reviewing segment composition, checking output labels against a sensitive-category list, or spot-checking a sample of outputs before activation; more sensitive or higher-scale uses may warrant closer review of segment composition or outcomes.
- Consider ongoing or periodic re-assessments, particularly after material model updates, prompt changes, or source-data changes. Note that drift can also occur passively through changes in the underlying data distribution, without any modification to the model or prompt.
- Consider what actions should be taken if outputs begin to cluster around sensitive traits or contexts.

DON'T

- ✗ **Don't assume that because nobody intentionally targeted a sensitive audience or topic, an AI-generated segment or audience will not be an effective proxy.**

4. Disclosure

Disclosure helps consumers, partners, and internal teams understand material AI-enabled processing and apply choices and rights to the relevant activities.

DO

Disclose AI involvement in data processing to consumers when it is material to how a consumer's information is used.

- Ask whether the AI use materially changes what a reasonable consumer would understand about how their information is collected, used, or shared. This may include what choices or rights are relevant, and how to exercise them. Depending on the context, uses of AI may call for disclosure if the AI system infers new attributes about consumers, constructs audience segments from broad or unstructured inputs, or enables autonomous data access that existing disclosures do not describe.
- Where disclosure is warranted, use plain language that materially and accurately describes what the AI does in functional terms, including the words "AI" or "machine learning" where they accurately describe the processing.
- Consider where consumer-facing disclosures should appear (e.g., privacy notices or other locations). The disclosure should be specific enough that a consumer who wants to exercise applicable rights, such as opting out of targeted advertising or profiling, can understand what processing their choice applies to.
- Advertiser and partner-facing disclosures are addressed in Section 1. Make sure legal, product, and technical teams agree on any external description before it is used.

DON'T

- ✗ Don't use vague or euphemistic language to describe meaningful AI use.
- ✗ Don't highlight AI-driven processing to business partners while failing to ensure that consumer-facing disclosures accurately describe that processing in plain language where disclosure is warranted.

5. Permissions and constraints

Permissions and constraints keep an agent from using data, tools, or execution pathways beyond the approved advertising workflow.

DO

Encode privacy and governance rules into agentic system permissions and constraints.

- Define the scope of authority for each agent or automation before deployment.
- Use allowlists, input validation, or similar controls for approved tools, datasets, APIs, and destinations, and guard against prompt or instruction manipulation that could redirect an agent outside approved pathways.
- Implement least-privilege and least-data access by default. Emerging bounded-execution architectures illustrate how constrained connectivity can support this objective in practice.²
- Consider requiring agents to declare a stated purpose or intent when requesting data access or proposing workflow changes, and avoid defining agentic authority or permitted purposes so broadly that the agent can access, infer, combine, or transfer data beyond what is reasonably needed for the approved advertising workflow.
- Where agentic workflows access third-party systems, platforms, APIs, or accounts, contracts should address whether agent access should be separately permissioned, identifiable, and limited to authorized pathways.

DON'T

- ✗ Don't give agents or AI tools open-ended access to data, tools, or execution pathways.
- ✗ Do not authorize agentic access to external data, onward transfer, or other use outside the approved scope unless such agentic use is explicitly documented and approved.

6. Choice and signal handling

Signal handling preserves consumer choices, such as opt-outs and consent, as data and instructions move through automated workflows and intermediaries.

DO

Ensure agentic workflows honor applicable consumer choice and privacy signals.³

- Confirm where in the workflow signals are received, processed, passed through, and enforced.
- Consider how applicable signals such as Global Privacy Control (GPC), Global Privacy Protocol (GPP) signals, or comparable preference signals are honored and preserved as data moves across intermediaries, orchestration layers, or partner systems.
- Where automated or agentic systems can generate, alter, or suppress choice signals (for example, consent or opt-out indicators), design those systems to preserve the integrity and provenance of consumer signals and take reasonable steps to prevent the creation or modification of signals in ways that misrepresent a consumer's actual choices.
- Treat signal-handling failures as governance failures requiring remediation, not merely engineering bugs (see also Section 7 below).
- Periodically test and document how agentic workflows respond to and honor applicable consumer choice and privacy signals.

DON'T

- ✘ Don't allow privacy signals to be dropped, ignored, or stranded as data moves across agentic workflows.
- ✘ Don't assume privacy signals are being properly honored without consistent testing of how agentic workflows respond to these signals.

7. Oversight and logs

Oversight lets members detect, investigate, and stop unexpected agentic activity before it risks becoming a broader privacy, security, or contractual problem.

DO

Maintain human escalation, kill-switch, and audit-log controls proportionate to the system's capabilities.

- Define escalation triggers for sensitive contexts, policy exceptions, or authority-limit breaches.
- Maintain a tested means to promptly suspend, constrain, or disable an agentic system's actions (i.e., a kill switch) when the system behaves unexpectedly or exceeds its defined scope. The mechanism should be capable of responding at the speed at which the system operates; for agentic workflows executing at programmatic speeds, this would typically mean automated triggers in addition to manual intervention.
- Keep logs sufficient to reconstruct what an agentic system did, what inputs and instructions led to its outputs, what signals or policies applied, and what human approvals occurred. Keep these logs proportionate to the system's risk and minimize the personal or sensitive data they capture, applying appropriate access controls and retention limits. Primarily, this is expected to involve logging information about the decision context, not model weights or proprietary logic. For execution-capable agents, consider assigning a unique, persistent identifier to each agent or agent instance, where technically feasible, and linking that identifier to logs showing the initiating user or customer, declared purpose, access pathway, applicable constraints, and resulting actions.
- Consider whether data being logged are sufficient to evaluate data lineage, chain of custody, or whether use remained within the authorized scope.
- Be prepared to take remediation steps proportional to any issues identified through oversight measures. For example, when monitoring or oversight surfaces a control failure or unexpected agentic activity, be prepared to (a) suspend or constrain the affected workflow; (b) preserve relevant logs; (c) notify affected partners or platforms where appropriate; (d) remediate the underlying control failure; and (e) document the issue and its resolution.

DON'T

- ✘ **Don't treat all AI tools in the same risk tier. Advisory and informational AI tools may not warrant the same level of oversight as execution-capable agentic AI tools, but still require human accountability for the outcomes they inform.**

8. Contracting

Contracting closes responsibility gaps among providers, customers, and platforms before a workflow is deployed, or before something goes wrong.

DO

Address responsibility for AI-enabled tools and workflows expressly in contracts, so that both providers and customers have a shared understanding of who is responsible for what.

- Members should consider how responsibility is allocated between the provider of an AI-enabled tool and the advertising customer that uses it. The line will not always be clean and will depend on each member's role, but contracts should address it rather than leave it assumed.
- Members that license or deploy third-party AI tools should conduct reasonable diligence on the provider's data sources, governance controls, security practices, and the permitted scope of data use before deployment, and revisit that diligence as capabilities change.
- Provider-side responsibilities typically include the system's design controls, documented use limitations, input and output constraints, and enabling appropriate transparency and auditability to the customer.
- Customer-side responsibilities typically include the inputs the customer provides, the activation decisions the customer makes based on AI outputs, and exercising available review and oversight opportunities.
- Some responsibilities will be shared or will need to be negotiated based on the specific relationship. For example, for consumer-facing disclosures, the provider may need to supply information about how its AI system works for accurate disclosures, but the customer is more likely to be responsible for making those disclosures to consumers. Signal handling across intermediaries and remediation if something goes wrong may also be shared responsibilities.
- Where natural-language audience creation or other generative capabilities are available, customer-side protections may include requiring that sensitive categories be excluded from the source taxonomies or datasets available to the tool.
- Vendor-side protections may include commitments from the customer not to intentionally target prohibited sensitive categories, even through proxy categories. Both approaches may be combined.
- For agentic workflows, consider how provider and customer responsibilities map to each party's operational role. For example, where an SSP or exchange hosts an agentic system that executes within its auction infrastructure, the SSP or exchange is likely better positioned to control what data the agent can access, log what the agent did, and enforce approved routing and destinations. The ad-tech provider whose agent runs within that infrastructure is typically better positioned to honor the privacy signals it receives, operating within its declared purpose, and not exceeding the scope of its approved access. Contracts between these parties should address notice, cooperation, and remediation steps if the system exceeds agreed constraints or mishandles data or signals.⁴

- Contracts should allocate responsibility for third-party platform access, including who ensures that agent activity occurs through authorized access methods and can be distinguished from human user activity.
- Contracts should clearly define what uses of data exposed during AI-enabled workflows are authorized and prohibited, including: retention, reuse, onward transfer; combination with other datasets; use of campaign performance data or audience inputs for model fine-tuning or general training; and uses outside the agreed scope.

DON'T

- ✘ Don't leave responsibility for AI-specific risks unaddressed in contracts or negotiations.

9. Accountability

Cross-functional accountability keeps AI advertising governance operationally enforceable, legally sound, and accurately represented to the market.

DO

Assign cross-functional accountability for agentic workflows. Consider how:

- Product, engineering, and security teams should be involved in technical enforcement, security review, and logging, including assessing risks such as unauthorized tool use, prompt or instruction manipulation, unintended data exposure, and uncontrolled onward sharing.
- Legal and privacy teams should help define disclosures, risk thresholds, and contracting positions.
- Commercial teams will communicate what the product does and does not deliver.
- Different stakeholders can determine when a workflow requires escalation rather than routine approval.

DON'T

- ✘ Don't silo AI advertising governance inside any single team or function.

C. Suggested implementation checklist

Members adopting AI in Network Advertising should consider whether they can answer “yes” to the following working questions. These are primarily intended for systems that meet the scope criteria above; for established optimization and modeling tools, existing NAI Framework compliance remains the baseline.

#	CONTROL	WORKING QUESTION
01	Inventory	Do we maintain an inventory of AI-enabled advertising tools, models, agents, and outputs, categorized by level of autonomy?
02	Segment review and activation	Do we review AI-generated segments and their activation context, including sensitive verticals, before use?
03	Testing	Do we conduct reasonable evaluation for proxy risk and drift, proportionate to the system’s capabilities?
04	Disclosure	Do our privacy notices and product materials describe AI uses that materially change how consumer data is processed?
05	Permissions	Do we define authority, approved access, and system constraints before an agentic system is deployed?
06	Signal handling	Do we preserve and honor applicable preference and privacy signals throughout agentic workflows, and protect the integrity of those signals?
07	Oversight	Do we maintain escalation triggers, a kill switch where appropriate, and auditable logs for higher-risk systems?
08	Contracting	Do our contracts and controls define authorized uses of workflow-exposed data and allocate AI-specific responsibilities among the relevant parties?
09	Accountability	Have we assigned cross-functional ownership, escalation authority, and review responsibilities for in-scope workflows?

D. Notes

1. See, e.g., IAB, [AI Governance and Risk Management Playbook](#) (Aug. 26, 2025); IAB, [AI Transparency and Disclosure Framework](#) (Jan. 15, 2026); IAB Tech Lab, [AAMP](#) (Agentic Advertising Management Protocols) (last updated Apr. 16, 2026); Rowena Lam, [“As AI Agents Transform Digital Advertising, Where’s the Privacy Architecture?”](#), Cynopsis (Mar. 9, 2026).
2. See, e.g., IAB Tech Lab, [Agentic Real Time Framework \(ARTF\)](#) standards materials (last updated Feb. 18, 2026) and *Agentic Real Time Framework Specification v1.0* (released Nov. 13, 2025; and subsequently renamed from the “Agentic RTB Framework”); [Ad Context Protocol \(AdCP\)](#), Getting Started / protocol documentation and reference implementation.
3. See, e.g., IAB Tech Lab, [Global Privacy Protocol \(GPP\)](#) (last updated Dec. 17, 2025); IAB Tech Lab, ARTF standards materials (last updated Feb. 18, 2026).
4. See, e.g., IAB Tech Lab, ARTF v1.0; Ad Context Protocol (AdCP), protocol documentation. The allocation of responsibilities may vary depending on the architecture adopted.