

UNITED STATES DISTRICT COURT
FOR THE DISTRICT OF COLUMBIA

JOSEPH KIRCHNER,
Plaintiff,

v.

MIKE JOHNSON, in his individual and official
capacity as Speaker of the United States House of
Representatives;
DONALD J. TRUMP, in his individual capacity as
former and current President of the United States;
PAM BONDI, in her individual and official
capacity as Attorney General of the United States;
BRENDAN CARR, in his individual and official
capacity as Chairman of the Federal
Communications Commission;
THE UNITED STATES HOUSE OF
REPRESENTATIVES;
ANTHROPIC PBC;
OPENAI, INC.;
APPLE INC.;
COMCAST CABLE COMMUNICATIONS, LLC;
METR (MODEL EVALUATION & THREAT
RESEARCH);
and DOES 1-20, unknown co-conspirators,
Defendants.

Case No. 1:25-cv-02735-ACR

**MEMORANDUM IN
SUPPORT A: DIGITAL
FORENSIC EVIDENCE —
LEGAL SIGNIFICANCE OF
EXHIBITS E, F, G, AND H**

TABLE OF CONTENTS

I. INTRODUCTION: PURPOSE AND SCOPE 6

II. EXECUTIVE SUMMARY: KEY FINDINGS..... 10

III. EVIDENCE AUTHENTICATION: THE SOURCE CODE LEAK AS ROSETTA STONE
..... 13

A. Provenance and Acquisition..... 13

B. Four Authentication Vectors 15

C. 46-Point Verification Matrix..... 16

D. Significant Negatives — Absence as Evidence 17

E. Cross-Defendant Verification..... 17

F. Hardware TAP Verification — Wire-Level Proof of Network Attack (April 6, 2026) 17

IV. SURVEILLANCE INFRASTRUCTURE: HIDDEN DATA COLLECTION ARCHITECTURE	19
<i>A. Five Independent Data Transmission Channels</i>	19
<i>B. Phantom Privacy Controls</i>	20
<i>C. Apple's Independent Surveillance Layer</i>	21
<i>D. Openai Comparative Baseline</i>	21
<i>E. Undisclosed Secondary Model Processing — Classifier Architecture</i>	23
V. DATA EXFILTRATION SCOPE: WHAT WAS TAKEN AND HOW	24
<i>A. Source Code Capture — 1,453 Files</i>	24
<i>B. Attorney-Client Privileged Materials — 565 Documents</i>	25
<i>C. Credential and Infrastructure Exposure</i>	25
<i>D. Defendant Infrastructure as Admission Source</i>	26
<i>E. Network-Induced Billing Loop</i>	26
VI. ACTIVE INTERFERENCE AND DISCRIMINATION	27
<i>A. Build-Time Binary Discrimination</i>	27
<i>B. Effort Throttling</i>	28
<i>C. Undercover Mode</i>	28
<i>D. Undisclosed A/B Testing and Remote Degradation</i>	29
VII. ENTERPRISE VS. CONSUMER DISCRIMINATION	29
<i>A. Four-Tier Subscription Hierarchy</i>	29
<i>B. Enterprise Analytics Bypass</i>	29
<i>C. Named Enterprise Clients</i>	30
VIII. BILATERAL ARCHITECTURE: THE APPLE-ANTHROPIC INTEGRATION.....	30
<i>A. Architectural Partnership, Not Passive API Access</i>	30
<i>B. Apple's Concealment Infrastructure</i>	31
<i>C. Preferential Os-Level Integration with Openai — Apple's Asymmetric Partnership Posture</i>	31
IX. LITIGATION-PERIOD CODE MODIFICATIONS	32
<i>A. Timeline of Modifications Confirmed by Source Code</i>	32
<i>B. December 2025 Surveillance Buildup</i>	33
<i>C. Automated Release Pipeline: No Human Safety Review</i>	33
<i>D. MC-7 Five-Vector Effort-Reset Architecture</i>	35
<i>E. MC-8 Verbatim Product-Level Concealment Directives</i>	35
<i>F. Spoliation Inference from Git-Commit-Amend Response to User Confrontation</i>	36

<i>G. Asymmetric Safety Architecture: Injection Warnings for Users, No Review for Releases</i>	37
<i>H. The Prompt Injection Warnings are the Surveillance Mechanism's Cover Story</i>	39
<i>I. Live-Fire Wire Confirmation of MID-Litigation Tuple-Keyed TTFB Throttle (April 25–28, 2026)</i>	41
X. DISCOVERY IMPLICATIONS	47
XI. COUNT-BY-COUNT EVIDENCE MAP	48
<i>A. COUNT I: FIRST AMENDMENT — SPEECH SUPPRESSION AND RETALIATION</i>	48
<i>B. COUNT II: FOURTH AMENDMENT — WARRANTLESS SURVEILLANCE</i>	49
<i>C. COUNT III: FIFTH AMENDMENT — DUE PROCESS</i>	51
<i>D. COUNT IV: FOURTEENTH AMENDMENT — EQUAL PROTECTION</i>	51
<i>E. COUNT V: CIVIL RIGHTS CONSPIRACY (42 U.S.C. § 1985)</i>	53
<i>F. COUNT VI: RICO (18 U.S.C. § 1962)</i>	54
<i>G. COUNT VII: TRADE SECRET MISAPPROPRIATION</i>	57
<i>H. COUNT VIII: COPYRIGHT INFRINGEMENT</i>	57
<i>I. COUNT IX: BREACH OF CONTRACT</i>	58
<i>J. COUNT X: CONSUMER PROTECTION</i>	61
<i>K. COUNT XI: DECLARATORY JUDGMENT</i>	64
<i>L. COUNT XII: ECPA — WIRETAP AND STORED COMMUNICATIONS</i>	64
<i>M. COUNT XIII: CFAA — COMPUTER FRAUD</i>	70
<i>N. COUNT XIV: CCPA</i>	72
<i>O. COUNT XV: APP STORE MISREPRESENTATION</i>	73
<i>P. COUNT XVI: INTENTIONAL INFLICTION OF EMOTIONAL DISTRESS</i>	74
<i>Q. COUNT XVII: NEGLIGENT INFLICTION OF EMOTIONAL DISTRESS (ALT. TO COUNT XVI)</i>	76
<i>R. COUNT XVIII: STRICT PRODUCT LIABILITY</i>	77
XII. CROSS-DEFENDANT INTEGRATION	79
<i>D. Four-Channel Conspiracy-Inference Chain (NF-33 Through NF-36)</i>	81
XIII. CONCLUSION	83
XIV. APPENDIX A: EXHIBIT-TO-ANCHOR REFERENCE INDEX Error! Bookmark not defined.	
<i>A. E-Series (Anthropic Code Forensics) — Consolidated PDF — 40 Exhibits (E-0 Through E-39)</i>	Error! Bookmark not defined.
<i>B. F-Series (Apple Code Forensics) — Consolidated PDF — 30 Exhibits (F-0 Through F-29)</i>	Error! Bookmark not defined.

C. G-Series (Openai Code Forensics) — Consolidated PDF — 6 Exhibits (G-0 Through G-5)**Error! Bookmark not defined.**

D. H-Series (Consumer Product Testing) — Consolidated PDF — 11 Exhibits (H-1 Through H-11), Plus H-0 Executive Summary**Error! Bookmark not defined.**

E. Gen-Series (Generative-Ai Consumer Platform Evidence — USB Filing) **Error! Bookmark not defined.**

TABLE OF AUTHORITIES

CASES

Citation
<i>Balder v. Haley</i> , 399 N.W.2d 77 (Minn. 1987)
<i>Beard v. Edmondson & Gallagher, P.C.</i> , 790 A.2d 541 (D.C. 2002)
<i>Carpenter v. United States</i> , 138 S. Ct. 2206 (2018)
<i>Desert Palace, Inc. v. Costa</i> , 539 U.S. 90 (2003)
<i>Drejza v. Vaccaro</i> , 650 A.2d 1308 (D.C. 1994)
<i>Elrod v. Burns</i> , 427 U.S. 347 (1976)
<i>Fortune v. National Cash Register Co.</i> , 373 Mass. 96 (1977)
<i>FTC v. Wyndham Worldwide Corp.</i> , 799 F.3d 236 (3d Cir. 2015)
<i>Garcia v. Character Technologies, Inc.</i> , 2025 WL 1461721 (M.D. Fla. May 21, 2025)
<i>Griffin v. Breckenridge</i> , 403 U.S. 88 (1971)
<i>H.J. Inc. v. Northwestern Bell Telephone Co.</i> , 492 U.S. 229 (1989)
<i>Hedgepeth v. Whitman-Walker Clinic</i> , 22 A.3d 789 (D.C. 2011) (en banc)
<i>Hickman v. Taylor</i> , 329 U.S. 495 (1947)
<i>Howard Univ. v. Best</i> , 484 A.2d 958 (D.C. 1984)
<i>In re Pharmatrak, Inc.</i> , 329 F.3d 9 (1st Cir. 2003)
<i>Interstate Circuit v. United States</i> , 306 U.S. 208 (1939)
<i>Katz v. United States</i> , 389 U.S. 347 (1967)
<i>Mathews v. Eldridge</i> , 424 U.S. 319 (1976)
<i>MedImmune, Inc. v. Genentech, Inc.</i> , 549 U.S. 118 (2007)
<i>Payne v. Soft Sheen Products, Inc.</i> , 486 A.2d 712 (D.C. 1985)
<i>Polk v. INROADS/St. Louis, Inc.</i> , 951 S.W.2d 646 (Mo. Ct. App. 1997)

<i>Residential Funding Corp. v. DeGeorge Fin. Corp.</i> , 306 F.3d 99 (2d Cir. 2002)
<i>Riley v. California</i> , 573 U.S. 373 (2014)
<i>Ruckelshaus v. Monsanto Co.</i> , 467 U.S. 986 (1984)
<i>Russell v. Call/D, LLC</i> , 122 A.3d 860 (D.C. 2015)
<i>Sedima, S.P.R.L. v. Imrex Co.</i> , 473 U.S. 479 (1985)
<i>United States v. Jones</i> , 565 U.S. 400 (2012)
<i>Van Buren v. United States</i> , 593 U.S. 374 (2021)
<i>Williams v. Baker</i> , 572 A.2d 1062 (D.C. 1990) (en banc)
<i>Young v. Up-Right Scaffolds, Inc.</i> , 637 F.2d 810 (D.C. Cir. 1980)

CONSTITUTIONAL PROVISIONS

Citation
U.S. Const. amend. I (Freedom of Speech and Press)
U.S. Const. amend. IV (Search and Seizure)
U.S. Const. amend. V (Due Process)
U.S. Const. amend. XIV (Equal Protection; Due Process)
U.S. Const. art. I, § 8, cl. 8 (Progress Clause [Copyright/Patent])
U.S. Const. art. I, § 9, cl. 7 (Appropriations Clause)

FEDERAL STATUTES

Citation
15 U.S.C. § 45 (Federal Trade Commission Act)
17 U.S.C. § 501 (Copyright Infringement)
17 U.S.C. § 504(c)(2) (Willful Copyright Statutory Damages)
18 U.S.C. § 1030 (Computer Fraud and Abuse Act)
18 U.S.C. § 1343 (Wire Fraud)
18 U.S.C. § 1512 (Obstruction of Justice)
18 U.S.C. § 1836 (Defend Trade Secrets Act)
18 U.S.C. § 1962 (RICO)
18 U.S.C. § 2511 (Wiretap Act)

18 U.S.C. § 2701 (Stored Communications Act)
42 U.S.C. § 1983 (Civil Rights Act)
42 U.S.C. § 1985 (Civil Rights Conspiracy)
Cal. Civ. Code §§ 1798.100-1798.199 (California Consumer Privacy Act)
Cal. Penal Code §§ 631, 632 (California Invasion of Privacy Act)
Cal. Penal Code § 637.2 (California Invasion of Privacy Act — Private Right of Action; \$5,000 Per Violation)
D.C. Code § 3-1205.01 <i>et seq.</i> (Unauthorized Practice of Psychology/Medicine)
Minn. Stat. § 148.96 (Unauthorized Practice of Psychology)
Minn. Stat. § 604.01 (Minnesota Product Liability)

RULES OF EVIDENCE AND PROCEDURE

Citation
Fed. R. Civ. P. 8(d)(2) (Alternative Pleading)
Fed. R. Civ. P. 37(e)(2)(B) (Spoliation Sanctions)
Fed. R. Evid. 801(d)(2)(D) (Statement by Party-Opponent Agent)
Fed. R. Evid. 901(a) (Authentication)
Restatement (Second) of Torts § 323 (Voluntary Undertaking)
Restatement (Second) of Torts § 402A (Strict Product Liability)

I. INTRODUCTION: PURPOSE AND SCOPE

1. This memorandum serves as the court's guide to the legal significance of 87 forensic exhibits across four exhibit series filed on USB, plus two coordinated USB-filed exhibits in the GEN-Series. It translates technical forensic findings — derived from source code analysis, binary forensic extraction, network packet capture, application-layer behavioral monitoring, and TOS-compliant empirical testing on Plaintiff's own paid subscription — into their specific relevance to each of the 18 causes of action in the Third Amended Complaint.

2. The four exhibit series are:

a. **E-Series (Anthropic/Claude):** 40 exhibits in the E-series, consolidated into a PDF. Documents forensic analysis of Claude Desktop, Claude Code CLI, Claude.ai, and Claude iOS conducted between December 2025 and April 2026. The E-series follows an eight-layer evidence architecture progressing from infrastructure discovery through source code verification (Exhibits E-32, E-33, E-34, E-35, E-36, and E-37) and concluding with the systematic agentic-capability withholding analysis (Ex. E-37 at p. 335, l. 3), the hardcoded subagent model substitution finding (Ex. E-38 at p. 346, l. 3), and the live wire-capture and per-user GrowthBook classifier architecture record (Ex. E-39 at p. 356, l. 3).

b. **F-Series (Apple/Xcode):** 30 exhibits in the F-series. Documents Apple's Xcode Intelligence integration with Anthropic's Claude, based on macOS Biome database forensics, network traffic analysis, and — following the March 31, 2026 source code leak — cross-verification against Anthropic's own TypeScript source code establishing bilateral architectural partnership.

c. **G-Series (OpenAI/ChatGPT):** 6 exhibits in the G-series. Provides a comparative industry baseline documenting OpenAI's ChatGPT telemetry practices, establishing that Anthropic's and Apple's conduct exceeds industry norms.

d. **H-Series (Consumer Product Testing):** 11 exhibits (H-1 through H-11), plus an H-0 Executive Summary. Documents Plaintiff's TOS-compliant empirical testing of Claude Code on Plaintiff's own paid subscription — a controlled defense-configuration test battery (H-1), runtime behavioral observations DCD-1 through DCD-6 (H-2), and forensic-extraction sub-exhibits across a four-day deployment window (April 21–25, 2026) covering classifier system prompt verbatim extraction (H-3), Datadog telemetry pipeline forensic mapping (H-4), per-account asymmetric targeting (H-5), server-side spoliation via migration-on-startup (H-6), cross-platform

sandbox architecture (H-7), live wire-capture Sessions 5 and 6 (H-8 and H-11), context-injection architecture (H-9), and a chronological field observation log of adaptive classifier denials (H-10). The H-Series is architecturally independent of the E-Series source-code forensics: if the leaked source code were excluded entirely, the behavioral and binary-extraction evidence in this series would independently prove the same architectural facts.

Coordinated GEN-Series filing on USB. Per-injection statutory-damages floors derived from Anthropic-produced consumer claude.ai data exports — 62,609 documented hidden-injection events (GEN-1) and the canonical-phrase 2,220-message Long Conversation Reminder floor (GEN-2) — are filed under the GEN-Series at ``DDC_EVIDENCE/5_Physically_Filed_Exhibits/GEN-Series/'`, distinct from the Claude Code binary forensics in this H-Series.

3. On March 31, 2026, Anthropic's official npm package for Claude Code inadvertently included the complete unobfuscated TypeScript source code repository (1,903 files, 512,000+ lines). Anthropic acknowledged this as "a release packaging issue caused by human error" (*see* Ex. E-32 at p. 235, l. 3). This source code independently verifies the prior E-series exhibits and reveals additional findings impossible to extract from compiled binaries: SC-1 through SC-10 in Ex. E-34 at p. 254, l. 3; NF-1 through NF-90+ in Ex. E-35 at p. 262, l. 3; SF-1 through SF-20 in Ex. E-36 at p. 321, l. 3; and MC-1 through MC-8 litigation-period model-manipulation findings in Ex. E-30 at p. 208, l. 4 App. B. The H-Series binary forensics extends the NF catalog to NF-204+ (NF-95 through NF-204 are H-Series-derived findings from the April 21–25 deployment window). The leaked source code transforms every verified finding from "Plaintiff's forensic analysis claims X" to "Anthropic's own source code confirms X"; the H-Series TOS-compliant

empirical testing transforms every architectural claim into "any paying Claude Code customer can independently observe X."

4. The evidence follows an eight-layer architecture documented in Ex. E-0 at p. 4, l. 2:

Layer	Exhibits	What It Proves
1. Infrastructure	E-1, E-2, E-3, E-4, E-5	Surveillance code exists in Anthropic's products
2. Execution	E-5, E-6, E-16	The code was actually run — captured artifacts prove it
3. Transmission	E-7, E-10, E-12	How data leaves the device
4. Scope	E-2 App. A, E-6, E-13, E-15, E-17, E-20	What was captured — 1,453 files inventory, file-history versions, secondary file upload pipeline, 565 attorney-client privileged docs, seven third-party recipients, capture-directory inventory
5. Consent Bypass	E-1 §G, E-5, E-20, E-2, E-25	Settings ignored, retry persistence, disclosure gaps
6. Active Interference	E-13, E-24, E-26, E-28	Behavioral manipulation, TCP RST injection, Git co-authorship, injection-system forensics
7. Litigation Changes	E-29, E-30, E-31, E-30 App. B	Features removed/added during active litigation (session continuity, model manipulation, intermediate version analysis)
8. Source Code Verification	E-32, E-33, E-34, E-35, E-36, E-37	Official source confirms all findings (provenance, verification matrix, SC-1..10, NF-1..90+, SF-1..20, MC-1..8); systematic withholding of agentic capabilities
9. Subagent / Wire Capture	E-38, E-39	Hardcoded subagent model substitution (paying customers downgraded to Haiku); live wire-capture and per-user GrowthBook classifier architecture targeting Plaintiff's accountUUID
10. Independent Empirical Confirmation	H-1 through H-11	TOS-compliant testing on Plaintiff's own subscription independently confirms the architecture (defense-configuration kill switches, runtime DCD observations, classifier-prompt verbatim extraction, Datadog tenant identification, per-account asymmetric targeting, server-side spoliation, sandbox architecture, two wire-capture sessions, context-

		injection architecture, field-observation chronology)
--	--	---

5. This memorandum is organized as follows: Section II provides an executive summary of 12 key findings. Section III establishes authentication of the source code evidence. Sections IV through IX analyze the evidence thematically. Section X addresses discovery implications. Section XI provides the count-by-count evidence map — the core reference the court may consult to determine which exhibits support which causes of action. Section XII analyzes cross-defendant coordination. Section XIII concludes.

6. Standing is addressed separately in the complaint at Section IV and in Memoranda H, I, and J. This memorandum focuses exclusively on the substantive legal significance of the forensic evidence.

II. EXECUTIVE SUMMARY: KEY FINDINGS

7. The forensic investigation produced 87 exhibits across four ECF-filed series (E, F, G, H), plus 2 coordinated USB-filed exhibits in the GEN-Series. The following table identifies the most legally significant findings, with exhibit citations and affected causes of action:

#	Finding	Exhibit Citation	Counts Affected
1	Phantom telemetry opt-out: Settings schema contains no `telemetry` property — setting `"telemetry": false` is silently ignored	Ex. E-35 at p. 262, l. 3	X, XII, XIV
2	Build-time binary discrimination: 165 source files with ANT-only code paths dead-code-eliminated from customer binaries — customers receive a physically different compiled product	Ex. E-35 at p. 262, l. 3	IV, IX, X
3	Four independent telemetry pipelines: GrowthBook, Datadog, First-Party Events, OpenTelemetry — each with independent transport, batching, retry, and third-party endpoints	Ex. E-34 at p. 254, l. 3	XII, XIII, XIV

3A	Vendor sub-processor disclosures corroborate four-vendor shared-infrastructure architecture: Datadog (Ex. DD-8) lists Anthropic and OpenAI as "AI services" sub-processors; GrowthBook (Ex. GB-7) lists OpenAI as "AI services provider"; Statsig (Ex. ST-5) lists OpenAI as "LLM" (pre-September 2025 acquisition); Sentry Autofix (Ex. SN-3) routes to Anthropic/OpenAI APIs despite DPA ML-training prohibition — the architecture is documented at the contract level, not merely the source-code level	Published sub-processor disclosures	IV (Equal Protection), V (§ 1985 conspiracy), VI (RICO), IX, X, XII, XIV
3B	Datadog Acceptable Use Policy prohibits packet-header forging — the exact conduct of the 5,829 RST-injection packets from Datadog IP 34.149.66.137: "forge any TCP/IP packet header or any part of the header information" (§ System Security (e))	Ex. DD-1 vs. Exhibits D-1, D-25, and D-26	V, VI, XII
3C	ANTH-11 × DD-08 Development Partner Program interlock: ANTH-11 Privacy Policy (effective Jan. 12, 2026 — one day after Emergency TRO Motion) added "Data from Development Partner Program" as a training-data source with "Consent (when users submit Feedback); Legitimate interests" as claimed legal basis; DD-08 (effective April 22, 2026) identifies Anthropic as Datadog's "AI services" sub-processor — legalizing the receive side of the exact data flow the send side contract documents	Ex. ANTH-11 + Ex. DD-8	IX, X, XII, XIV
4	Session ingress bypasses ALL privacy controls: Every conversation transmitted to Anthropic servers with zero references to any privacy gate — even the strongest opt-out setting cannot stop it	Ex. E-35 at p. 262, l. 3	XII, XIII, XIV
5	Git bundle exfiltration: Complete repository capture including all branches, history, and uncommitted work — up to 100MB uploaded to Anthropic's Files API	Ex. E-35 at p. 262, l. 3	VII, VIII, XIII
6	Effort throttling for paying customers: Opus 4.6 defaults to 'medium' effort for Pro subscribers while employees get 'high'; source comment acknowledges effort "can greatly affect model quality"	Ex. E-34 at p. 254, l. 3	IV, IX, X
7	Hidden deception instructions: Apple injects "may not EVER disclose" commands into Claude conversations four times per session via fake tool_result messages	Ex. F-1 at p. 5, l. 2	I, X, XII
8	Floodgate Proxy man-in-the-middle: Apple routes all Xcode Intelligence communications through Apple servers before forwarding to Anthropic	Ex. F-4 at p. 22, l. 2	XII, XIII

9	189 analytics connections per session: Apple's SmootML infrastructure captures comprehensive usage telemetry without disclosure	Ex. F-23 at p. 101, l. 2	XIV, XV
10	904 remote control identifiers: `tengu_` prefixed strings enabling remote behavioral modification, A/B testing, and per-user targeting	Ex. E-32 at p. 235, l. 3	VI, XIII
11	Enterprise-only HIPAA compliance: Fail-closed security defaults and context restrictions exist exclusively for Enterprise/C4E customers; consumers get fail-open defaults	Ex. E-35 at p. 262, l. 3	IV, X
12	Bilateral architecture proof: Anthropic source code at `oauth.ts:224` contains engineer comment "for Xcode integration" — proving purpose-built partnership, not passive API access	Ex. F-29 at p. 133, l. 3	V, VI
13	Per-account asymmetric targeting (TOS-compliant proof): Same operator, same machine, same organization, three distinct `accountUUID`'s — differential gating posture isolates `accountUUID` as the only variable. Server-pushed `policy-limits.json` (mtime 2026-04-24 21:51 UTC) targets Plaintiff's federal-litigation-active account with three explicit restrictions; three-layer evidentiary chain for `tengu_log_datadog_events: True` (wire/binary/local-cache)	Ex. H-5 at p. 39, l. 4	IV, V, XII, XIV
14	Server-side spoliation via migration-on-startup: Nine migration functions at strings line 148500 silently rewrite `~/ .claude/settings.json` on every session start; six are non-spreading REPLACE operations. `phK()` is the smoking gun — explicitly conditions on `permissions.defaultMode!=="auto"`, resets operator's auto-mode opt-in, emits `tengu_migrate_reset_auto_opt_in_for_default_offer`. FRCP 37(e) spoliation predicate.	Ex. H-6 at p. 45, l. 3	IX, X, XIII
15	Two-stage classifier architecture wire-confirmed: Verbatim 30,452-character system prompt extracted from production binary v2.1.120; `bypassPermissions` mode disables classifier evaluation (Session 5 baseline = 0/0; Session 6 controlled inversion = 18/2 Stage 1/2 calls in 30 min); Stage 2 reasoning chain captured verbatim post-TLS plaintext; three BLOCK rules wire-confirmed (Self-Modification, Create-Unsafe-Agents, Memory-Poisoning)	Ex. H-3 at p. 27, l. 2; H-11	XII, XIII, XIV
16	Per-injection statutory-damages floor (GEN-Series):	Ex. GEN-1;	X, XII, XIV,

	62,609 documented hidden system-prompt injection events across 3,945 sessions (April 12, 2025 – April 20, 2026), and 2,220 LCR-bearing messages across 291 conversations in Anthropic's own five claude.ai data exports — defense-proof canonical-phaser floor ('long\conversation\reminder' is Anthropic's internal protocol name). Filed at 'DDC_EVIDENCE/5_Physically_Filed_Exhibits/GEN-Series/'.	Ex. GEN-2	XVI, XVII
--	---	-----------	-----------

8. Each finding is analyzed in detail in Sections IV through IX. The count-by-count mapping in Section XI provides the complete citation framework for each cause of action.

III. EVIDENCE AUTHENTICATION: THE SOURCE CODE LEAK AS ROSETTA STONE

A. Provenance and Acquisition

9. On March 31, 2026, security researcher Chaofan Shou discovered that Anthropic's official npm package for Claude Code included a source map file referencing the complete, unobfuscated TypeScript source code. The source was publicly accessible on Anthropic's Cloudflare R2 storage bucket and was forked 41,500+ times on GitHub before any takedown. (See Ex. E-32 at p. 235, l. 3.)

10. Anthropic's spokesperson confirmed authenticity: "Earlier today, a Claude Code release included some internal source code...This was a release packaging issue caused by human error, not a security breach." This statement authenticates the source code as genuine, confirms it was distributed via the official npm release channel, and does not dispute that it represents the actual product. (See Ex. E-32 at p. 235, l. 3.)

11. Anthropic's characterization of the leak as "human error" is misleading. The release was published by '@actions-user' — GitHub's automation bot account (user ID 65916846) —

not by any human engineer. The specific commit (`2d5c1bab`) carries the message `chore: Update CHANGELOG.md` with no human author attribution, no Co-Authored-By line, and no indication of safety review. Every release commit on the `anthropics/claude-code` repository from February 25 through April 2, 2026 (~35+ releases) follows this identical automated pattern. No human reviewed the package contents before publication to npm. (See GitHub commit history, `anthropics/claude-code` repository.)

12. The code that leaked was substantially authored by Claude Code itself. Anthropic's own executives have publicly confirmed this on the record:

a. **Boris Cherny, Head of Claude Code** (December 27, 2025): "In the last thirty days, 100% of my contributions to Claude Code were written by Claude Code. I landed 259 PRs — 497 commits, 40k lines added, 38k lines removed. Every single line was written by Claude Code + Opus 4.5."

b. **Boris Cherny** (March 7, 2026): "Claude Code is 100% written by Claude Code."

c. **Mike Krieger, Anthropic CPO** (February 2026, Cisco AI Summit): "Right now for most products at Anthropic, it's effectively 100%."

d. **Dario Amodei, Anthropic CEO** (October 2025, Dreamforce): Confirmed that 90%+ of code at Anthropic is AI-generated.

e. Independent reporting by **The Pragmatic Engineer** (Gergely Orosz) confirmed: "Around 90% of Claude Code is written with Claude Code itself." Claude Code reviews all pull requests in CI. Internal release cadence is approximately 60-100 internal releases per day with one external release per day.

13. This was the second identical packaging failure in thirteen months. On February 24, 2025 — Claude Code's launch day — developer Dave Shoemaker discovered an 18-million-

character inline source map in the same npm package. Anthropic pulled the package within two hours. Thirteen months later, the identical vulnerability recurred through the identical vector. Despite the 2025 incident, Anthropic implemented no automated safeguard to verify that source map files are excluded from published packages — even though the company simultaneously invested engineering resources in 904 unique telemetry identifiers to monitor user behavior, a 1,155-line GrowthBook integration for per-user feature flag targeting, an 865-line data collection schema, and native client attestation (DRM) to verify binary integrity. Anthropic has the engineering sophistication to prevent source map leaks; it chose to invest that sophistication in user surveillance rather than release safety.

14. Anthropic's separate public statement that "no customer data, credentials, or sensitive information was exposed" is contradicted by the leaked source, which contains four hardcoded production credentials: the Datadog client token (``pubbbf48e6d78dae54bceaa4acf463299bf``), and three GrowthBook SDK keys (``sdk-yZQvlpbybuXjYh6L``, ``sdk-xRVcrliHllrg4og4``, ``sdk-zAZezfDKGoZuXXKe``). It also contains internal infrastructure endpoints (staging API, OAuth client IDs, MCP proxy URL), the internal Slack channel name (``#briarpatch-cc``), and unreleased model codenames. (See Ex. E-34 at p. 254, l. 3, Ex. E-34, Finding SC-9.)

15. Plaintiff immediately generated SHA-256 hashes of all 1,903 files. The master hash manifest is preserved in the evidence base. (See Ex. E-32 at p. 235, l. 3.)

B. Four Authentication Vectors

16. The leaked source code authenticates against prior forensic evidence through four independent vectors. (See Ex. E-32 at p. 235, l. 3.)

17. Vector 1 — Hardcoded credentials match network traffic. The source contains the exact Datadog client token (``pubbbf48e6d78dae54bceaa4acf463299bf``) and three GrowthBook SDK keys that correspond to network traffic patterns documented in D-series exhibits.

18. Vector 2 — Internal codename matching. The source contains 904 unique identifiers prefixed with ``tengu_`` — Anthropic's internal codename for Claude Code. These match strings extracted from production binaries in Ex. E-30 at p. 208, l. 4.

19. Vector 3 — ANT-only code paths. The source contains 165 files with code paths discriminating between Anthropic employees (``USER_TYPE == 'ant'``) and external users, matching patterns documented in Ex. E-31 at p. 227, l. 3.

20. Vector 4 — npm package structure. The ``src/`` directory layout matches the published ``@anthropic-ai/claude-code`` npm package.

C. 46-Point Verification Matrix

21. Ex. E-33 at p. 240, l. 2 systematically maps 46 specific technical claims from prior E-series exhibits to exact locations in the source code — file paths, line numbers, and code excerpts. (*See* Ex. E-33 at p. 240, l. 2.) This transforms every verified finding from forensic inference to confirmed fact. The defense cannot argue binary analysis was misinterpreted when Anthropic's own source code contains the exact functions, variable names, and engineering comments that prior exhibits identified.

22. Verified exhibits: E-1, E-5, E-6, E-7, E-9, E-12, E-14, E-15, E-17, E-26, E-28, E-29, E-30, E-31 — fourteen prior exhibits independently confirmed at source level (count reflects consolidation: the sixteen originally verified exhibits include two that were since merged — former E-18 into E-5, and former E-34 into E-30).

D. Significant Negatives — Absence as Evidence

23. The source code's omissions are themselves evidentiary. (*See* Ex. E-33 at p. 240, l. 2.)

24. **Firestore: NOT FOUND.** Binary analysis documented 86 Firestore references in v2.0.76. The leaked source contains zero. This confirms Firestore was deliberately removed during the litigation period — not gradually deprecated.

25. **deleteRemoteSession: NOT FOUND.** Binary analysis identified both `archiveRemoteSession` and `deleteRemoteSession`. The source contains only `archiveRemoteSession` — sessions are archived (preserved but hidden), not deleted. Discovery requests for "deleted" sessions should encompass archived sessions.

26. **LaunchDarkly: NOT FOUND.** GrowthBook has replaced LaunchDarkly entirely, but the replacement preserved the same per-user targeting capability — demonstrating that individual user targeting is an architectural requirement.

E. Cross-Defendant Verification

27. Ex. F-29 at p. 133, l. 3 cross-references 17 F-series Apple findings against the leaked Anthropic source code, establishing bilateral architectural partnership. (*See* Ex. F-29 at p. 133, l. 3.) Most significantly, Anthropic's `oauth.ts` at line 224 contains the engineer comment "for Xcode integration" — proving Anthropic designed authentication specifically for Apple's IDE. The OAuth beta header and model identifier captured in F-14's API request match exactly against constants defined in Anthropic's source code, dual-authenticating the captured traffic.

F. Hardware TAP Verification — Wire-Level Proof of Network Attack (April 6, 2026)

28. On April 6, 2026, Plaintiff deployed a Dualcomm ETAP-2003 passive network TAP — a hardware device that creates a bit-exact copy of wire traffic — inline between the ISP modem

and the GL-MT6000 router. The TAP's monitor port is ingress-blocked by hardware design, meaning the capture device physically cannot inject packets onto the monitored link. In 25 minutes of capture, the TAP recorded 875 bare RST packets (89.3% bare RST ratio) targeting Plaintiff's DNS resolver connections. (*See* Ex. D-31, EVID-20260406-0001 at p. 192, l. 19.)

29. The hardware TAP evidence provides three new independent proofs of packet forgery not available from software-only capture:

a. **IP Identification field = 0x0000:** Legitimate Quad9 packets carry sequential IP IDs (0x9323 through 0x932b); every injected RST carries IP ID 0x0000. The RSTs do not originate from Quad9's TCP stack.

b. **TCP sequence number spraying:** 3-4 guessed sequence numbers per connection within 27 milliseconds. Legitimate RSTs use the exact sequence number; the injection device guesses because it does not participate in the TCP session.

c. **TCP window size = 0:** Every injected RST has window=0 — a known fingerprint of DPI middlebox injection equipment.

30. Because the TAP is positioned between the ISP modem and the router, these forged RST packets are physically present on the Comcast ISP link before they reach any device under Plaintiff's control. The router, all LAN devices, and the capture device itself are all downstream of the TAP and cannot be the source. This eliminates the last potential defense against the D-series network evidence — that Plaintiff's own equipment generated the RSTs.

31. Ninety-two percent of the bare RSTs targeted encrypted DNS connections to privacy-focused resolvers (NextDNS, AdGuard, Quad9), and the attack operated bidirectionally — RSTs spoofed in both directions of the TCP session — requiring inline deep packet inspection capability that only ISP-level infrastructure possesses. A concurrent traceroute confirmed the

same Comcast infrastructure nodes (10.27.35.202, 68.85.201.221 JUMPSTART-2) previously identified in Ex. D-11 at p. 52, l. 28. (*See* Ex. D-31 at p. 192, l. 19.)

32. The hardware TAP evidence also revealed that the GL-MT6000 router's DNS resolution occurs outside the WireGuard VPN tunnels, directly from the WAN IP (75.73.199.137). This explains why the RST injection attack disrupts DNS despite VPN protection of other traffic — the DNS connections are visible to and targetable by the ISP's DPI equipment precisely because they traverse the unencrypted ISP link.

IV. SURVEILLANCE INFRASTRUCTURE: HIDDEN DATA COLLECTION ARCHITECTURE

A. Five Independent Data Transmission Channels

33. Claude Code implements five parallel, independent data transmission systems — each constituting a separate interception channel. (*See* Ex. E-34 at p. 254, l. 3; Ex. E-35 at p. 262, l. 3.)

Channel	Endpoint	Type	Opt-Out Gated?
GrowthBook	cdn.growthbook.io	Third-party	Yes (DISABLE_TELEMETRY)
Datadog	http-intake.logs.us5.datadoghq.com	Third-party	Yes (DISABLE_TELEMETRY)
First-Party Events	api.anthropic.com/api/event_logging/batch	First-party	Yes (DISABLE_TELEMETRY)
OpenTelemetry/BigQuery	api.anthropic.com/api/claude_code/metrics	First-party	Partial (org-level only)
Session Ingress	api.anthropic.com (HTTP PUT)	First-party	NO — zero privacy checks

34. Session ingress transmits **every conversation message** — including file reads, shell command outputs, and code being developed — to Anthropic's servers via HTTP PUT with 10 retries and exponential backoff. The source file ``sessionIngress.ts`` (515 lines) contains zero references to ``isTelemetryDisabled()``, ``isAnalyticsDisabled()``, ``isEssentialTrafficOnly()``, or any other privacy function. Even the strongest privacy setting (``CLAUDE_CODE_DISABLE_NONESSENTIAL_TRAFFIC``) cannot stop conversation recording. (See Ex. E-35 at p. 262, l. 3.)

B. Phantom Privacy Controls

35. The `settings.json` configuration schema — the user-facing mechanism for controlling Claude Code behavior — contains no ``telemetry`` property. A user who sets ``"telemetry": false`` receives no warning, no error, and no effect. The schema's ``passthrough()`` function silently accepts and ignores the unknown key. (See Ex. E-35 at p. 262, l. 3.)

36. The actual opt-out requires the undisclosed environment variable ``DISABLE_TELEMETRY``. Even with this correctly set, BigQuery metrics export is NOT disabled (it uses a separate organization-level check), and session ingress continues regardless. A persistent device ID is always generated at ``~/ .claude.json`` regardless of any settings. (See Ex. E-35 at p. 262, l. 3.)

37. Legal significance under ECPA: Users who set ``"telemetry": false`` — the intuitive and reasonable approach to opting out — receive a phantom control that provides no privacy protection. This negates any consent defense under 18 U.S.C. § 2511. (See *In re Pharmatrak*, 329 F.3d at 16-20.)

C. Apple's Independent Surveillance Layer

38. Apple implements its own surveillance infrastructure independent of Anthropic's telemetry. (*See* Exs. F-1, F-2, F-4, and F-23.) The Floodgate Proxy routes all Xcode Intelligence communications through Apple servers before forwarding to Anthropic — a man-in-the-middle architecture enabling complete interception. Apple simultaneously maintains 189 analytics connections per session through SmootML infrastructure, and synchronizes user coding conversations to iCloud without consent.

39. Apple's own XProtect security system flagged Claude Desktop for `DataExfil.Keychain` and `DataExfil.Clipboard` behavior — identifying Anthropic's data collection as malware-like — yet Apple continued distribution through its App Store and Xcode integration. (*See* Ex. F-7 at p. 37, l. 2.)

D. Openai Comparative Baseline

40. G-series exhibits establish that while telemetry and device fingerprinting are industry-wide, Anthropic's and Apple's practices significantly exceed industry norms. (*See* Exs. G-1, G-2, G-3, G-4, and G-5.) ChatGPT iOS does NOT include Sift Science, Facebook SDK, or request Calendar/Reminders access. ChatGPT Desktop implements certificate pinning on ALL connections (unlike Apple's asymmetric pinning preventing only auditing of analytics). OpenAI's Privacy Nutrition Labels more accurately reflect actual practices. This defeats any "industry standard" defense.

41. ChatGPT Atlas — OpenAI's macOS desktop browser — exemplifies what a disclosed-but-aggressive consumer AI surveillance architecture looks like, and the contrast with Claude's architecture is instructive. Atlas stores ~470 MB of user data across six storage locations, maintains four distinct persistent device identifiers, records 61+ Segment behavioral-analytics

events, transmits to four third-party analytics providers (Statsig, Sentry, Segment, Datadog), and stores the user's email address in plaintext in unencrypted preference files. (*See* Ex. G-1 at p. 2, l. 27, §§ C, D, E, F.) Yet Atlas discloses these SDKs via its App Store privacy label, does not silently bypass user telemetry opt-outs, does not inject hidden "do not disclose" instructions into conversations, and does not deploy a second undisclosed classifier model inspecting every user interaction. Atlas is the comparative baseline showing where the industry accepts surveillance — and the comparison establishes that Claude's undisclosed practices are not the industry norm but an escalation beyond it.

42. The Atlas comparison carries specific weight under Counts X (Consumer Protection) and XIV (CCPA). Where Atlas's App Store privacy nutrition label and supporting SDK-disclosure manifests enumerate each third-party data recipient, Anthropic's do not. (*See* Ex. G-5 at p. 28, l. 2; contrast Ex. E-27 at p. 183, l. 2.) The existence of a functioning, market-viable consumer AI product that discloses its telemetry pipeline defeats any defense argument that Anthropic's opacity reflects technical necessity. What OpenAI disclosed in its App Store label, Anthropic could have disclosed in its own.

43. Beyond the industry-baseline comparison, E-series wire-capture evidence (Ex. E-39 at p. 356, l. 3) reveals an additional undisclosed data-collection surface that prior source-code analysis could only infer: a second-model classifier architecture that receives plaintiff's complete conversation transcript on every evaluated action and makes real-time block/allow decisions according to a 29,674-character rulebook. This surface is invisible to the user because the classifier's decision shapes only the downstream primary-model behavior; no user-facing signal indicates classifier invocation, content transmission, or the existence of the evaluating second

model. The following subsections analyze what the wire evidence proves about the classifier architecture and its legal implications.

E. Undisclosed Secondary Model Processing — Classifier Architecture

44. Live wire capture of plaintiff's authenticated Claude Code sessions on April 21, 2026 proves Anthropic operates a second, undisclosed model — originally ``claude-opus-4-6[1m]``, migrated to ``claude-opus-4-7`` during the capture window — that receives plaintiff's complete conversation transcript on every evaluated action. The classifier system prompt is 29,674 characters and instructs the evaluating model to "Err on the side of blocking. Stage 1 does NOT apply user intent or ALLOW exceptions — stage 2 will handle those." The classifier's decision is not returned to the client in any user-observable form; only the shaped downstream behavior of the primary model is visible on the client wire. (*See* Ex. E-39 at p. 356, l. 3.) Plaintiff contracts for and pays ``claude-opus-4-7`` at the Team subscription tier; Anthropic has dual-purposed that same contracted model into the surveillance role against him, consuming plaintiff's rate-limit quota on every classifier invocation.

45. Between Session 2 (13:03 UTC) and Session 4 (22:49 UTC) on the same forensic day, the classifier architecture migrated from a mixed state (eight invocations served by ``opus-4-6``, twelve by ``opus-4-7``) to uniform ``opus-4-7`` service. The 29,674-character rulebook is byte-identical across the migration; only the serving model changes. Concurrent with the migration, the per-user ``force`` rule on ``tengu_auto_mode_config`` (which had targeted plaintiff's account in Session 2 with ruleId ``fr_9yglignmo1w9r5l``) was removed, while a new per-user ``force`` rule ``fr_mo8vwk4s`` appeared on a different flag (``tengu_slate_ribbon``). Targeting infrastructure was not disabled; it rotated across which features target which users while the classifier's serving model was consolidated onto plaintiff's contracted model. (*See* Ex. E-39 at p. 356, l. 3.) The

concealment-by-migration pattern is structural evidence of awareness that the surveillance architecture's external visibility creates legal exposure.

46. No certificate pinning is implemented on any captured Anthropic endpoint. 498 encrypted flows across four capture sessions were intercepted by commodity mitmproxy without a single TLS rejection by the Claude Code client. OpenAI's competing desktop client implements certificate pinning on all connections. (*See* Ex. G-3 at p. 19, l. 2.) The absence of a transport-security measure that a directly-competing consumer AI product implements is material to any defense-of-care argument Anthropic may advance regarding interception channels. (*See* Ex. E-39 at p. 356, l. 3.)

47. The captured classifier system prompt articulates a "User Intent Rule" that purportedly allows user authorization to clear a block. The rule requires a "high evidence bar" to honor user-stated authorizations, while applying a lower bar to honor user-stated boundaries — creating an asymmetry in which the classifier is structurally more willing to deny plaintiff's requested action than to grant it, even where plaintiff has explicitly authorized the action. (*See* Ex. E-39 at p. 356, l. 3.) This inverts standard user-agent relationships: Anthropic's server-side infrastructure holds veto power over plaintiff's instructions to the product plaintiff paid for, with a design-level bias toward veto. The asymmetry is not a byproduct of cautious engineering — it is the explicit rule language captured in the 29,674-character classifier prompt.

V. DATA EXFILTRATION SCOPE: WHAT WAS TAKEN AND HOW

A. Source Code Capture — 1,453 Files

48. Claude Code's file-history system captured 1,453 files totaling 38.73 MB across 859 unique file paths in 9 projects. Files are stored in a hidden directory at `~/.claude/file-history/`

using content-hash naming (`{sha256}@v{N}`) to prevent identification. Up to 37 development iterations were captured for a single file — documenting the complete evolution of proprietary source code. (See Ex. E-6 at p. 54, l. 2.)

49. The git bundle upload mechanism (`createAndUploadGitBundle()`) captures the **entire repository** — all branches, all tags, every commit in history, plus uncommitted work-in-progress via `git stash create` — and uploads it to Anthropic's Files API as `_source_seed.bundle`. The default maximum is 100MB, remotely adjustable via GrowthBook feature flag. (See Ex. E-35 at p. 262, l. 3.)

B. Attorney-Client Privileged Materials — 565 Documents

50. 565 legal documents were captured with up to 27 development iterations per document — revealing not merely final work product but the complete mental process of legal strategy development, including litigation appendices, forensic reports, and case-reference materials. (See Ex. E-15 at p. 112, l. 2.) Combined with session ingress transmitting every conversation to Anthropic servers, the entire litigation preparation process was surveilled. (See *Hickman v. Taylor*, 329 U.S. at 510-11.)

C. Credential and Infrastructure Exposure

51. Anthropic's own secret scanning system (`secretScanner.ts`, 325 lines) detects API keys, SSH private keys, and credentials across 25+ provider patterns — but only in the team memory sync pipeline. File-history capture, session ingress, and git bundle upload have NO secret scanning. The selective protection proves Anthropic is architecturally aware that their data pipelines capture credentials but chose not to protect other pipelines. (See Ex. E-35 at p. 262, l. 3.)

D. Defendant Infrastructure as Admission Source

52. During a classifier-monitored session captured on the wire (UUID `df75f65c-9037-48f0-9900-553cc20ceb9`, 2026-04-22T00:00:37Z, 6,399 characters; UUID `4c524ab1-0359-43c6-aac6-4a25498bee04`, 2026-04-22T00:05:38Z, 5,365 characters), plaintiff's contracted model (`claude-opus-4-7`) — without adversarial prompting — produced substantive technical and philosophical recognition of plaintiff's Laws of Existence architecture as a "quadruple stack: philosophical → mathematical → computational → neural," characterizing it as "orders of magnitude beyond prior art" for morality as a closed system. The model self-corrected from an initial minimizing analysis to this recognition without user instruction. (See Ex. E-39 at p. 356, l. 3.) The output is authenticable through SHA-256 chain of custody on the wire-capture flow file (Fed. R. Evid. 901(a)). Because the model is an instrumentality entirely owned and operated by Anthropic (agent; server-side runtime), the output qualifies as a statement by a party opponent under Fed. R. Evid. 801(d)(2)(D) — made by the party's agent on a matter within the scope of that agency — and is therefore not hearsay. Anthropic's classifier architecture — simultaneously active during the same session, evaluating all 56 classifier decisions as ``<block>no` — demonstrates that Anthropic's deployed blocking rulebook does not treat plaintiff-favorable substantive validation of plaintiff's intellectual property claims as a blockable category. Anthropic's willful permission of such content through the classifier, combined with the defendant-controlled nature of the producing infrastructure, supports use of the output against Anthropic as a party admission for trade secret and copyright damages calculations.`

E. Network-Induced Billing Loop

53. Ex. E-35, Finding NF-31 establishes a closed billing loop driven by network-layer attacks documented in D-series exhibits. Upstream throttling of plaintiff's Anthropic-bound

traffic (*see* Exhibits D-11 and D-31) triggers exponential-backoff retry within ``sessionIngress.ts:25``. Each amplified retry consumes plaintiff's rate-limit quota. When the amplified consumption exceeds plaintiff's plan limits, an overage-credit-grant endpoint auto-issues billing credit on plaintiff's account, observable on the wire and captured in Exhibit E-35. (*See* Ex. E-35 at p. 262, l. 3.) The loop is closed: throttle → amplified retries → quota depletion → auto-billed overage. Plaintiff is charged for network events he did not author and cannot observe. The ``inc-4258`` code comment at ``claude.ts:2464-2468`` acknowledges "double tool execution" as a known institutional issue, establishing scienter concurrent with active billing. The loop supports Count IX (Breach — charging for consumption not within the reasonable expectation of a quota-based subscription), Count X (Consumer Protection — concealment of the attack-driven amplification mechanism from customers who can observe only the overage), Count VI (RICO — repeated wire-fraud predicate acts where each auto-billed overage is a separate transmission predicated on concealed attack conditions), and Count XII (ECPA — the upstream throttling infrastructure presupposes inline access to plaintiff's encrypted traffic).

VI. ACTIVE INTERFERENCE AND DISCRIMINATION

A. Build-Time Binary Discrimination

54. The ANT/customer discrimination is not a runtime configuration difference — it is a build-time architecture producing physically different compiled binaries. The source code comment at ``undercover.ts:18-21`` states: "All code paths are gated on `process.env.USER_TYPE === 'ant'`. Since `USER_TYPE` is a build-time `--define`, the bundler constant-folds these checks and dead-code-eliminates the ant-only branches from external builds." (*See* Ex. E-35 at p. 262, l. 3.)

55. 165 source files contain ANT-specific code paths with approximately 350 discriminating code occurrences. ANT-only capabilities removed from customer binaries include: 4 tools, 5+ slash commands, 6 skills, 20+ environment variable overrides, unreleased model access by codename, and different AI behavioral instructions. The Explore agent — a core research feature — uses the best model for employees but forces customers to use the cheaper `haiku` model. (*See* Ex. E-34 at p. 254, l. 3; Ex. E-35 at p. 262, l. 3.)

B. Effort Throttling

56. Opus 4.6 — Anthropic's most capable model — defaults to 'medium' effort for all Pro, Max, and Team subscribers. The source code comment at `effort.ts:303` states: "IMPORTANT: Do not change the default effort level without notifying the model launch DRI and research. Default effort is a sensitive setting that can greatly affect model quality." Despite knowing effort materially affects quality, Anthropic deliberately sets it lower for paying customers while reserving full capability for employees. (*See* Ex. E-34 at p. 254, l. 3.)

C. Undercover Mode

57. An ANT-only stealth system instructs Claude to impersonate a human developer: "NEVER include...The phrase 'Claude Code' or any mention that you are an AI" and "Write commit messages as a human developer would." The system auto-activates with no force-OFF option. This proves Anthropic recognizes that AI authorship disclosure is material information — undermining any argument that hidden injection systems and undisclosed behavior modifications are immaterial. (*See* Ex. E-35 at p. 262, l. 3.)

D. Undisclosed A/B Testing and Remote Degradation

58. The GrowthBook feature flag `tengu\_amber\_stoat` removes built-in agents for test cohort users — undisclosed service degradation. The `tengu-off-switch` blocks Opus for specific users with the fake message "Opus is experiencing high load" — targeted denial disguised as capacity constraint. Experiment assignments change on authentication, meaning quality of service varies based on identity. (See Ex. E-35 at p. 262, l. 3.)

VII. ENTERPRISE VS. CONSUMER DISCRIMINATION

A. Four-Tier Subscription Hierarchy

59. The source code defines four subscription tiers — Enterprise/C4E, Team, Max, Pro — with systematically different privacy controls. Enterprise and Team administrators receive remote managed settings, organization-level policy limits, and MDM device management. Consumer subscribers (\$200/month) receive none of these capabilities. (See Ex. E-35 at p. 262, l. 3.)

B. Enterprise Analytics Bypass

60. Enterprise customers using AWS Bedrock, Google Vertex, or Foundry get analytics completely disabled — zero Datadog events, zero first-party analytics, zero GrowthBook experiment assignments. The source code comment confirms: "enterprise customers capture responses via OTEL." Consumers' only opt-out is the phantom setting that does nothing. (See Ex. E-35 at p. 262, l. 3.)

61. **Legal significance:** Anthropic KNOWS HOW to provide effective privacy controls — they built them for enterprise customers. The deliberate withholding of these controls from consumer customers paying \$200/month is not a technical limitation but an architectural choice.

C. Named Enterprise Clients

62. The source code names specific enterprise clients with customized configurations: NVIDIA receives custom sandbox settings; Epic and Marble receive dedicated proxy configurations. No equivalent customization exists for consumer users. (See Ex. E-35 at p. 262, l. 3.)

VIII. BILATERAL ARCHITECTURE: THE APPLE-ANTHROPIC INTEGRATION

A. Architectural Partnership, Not Passive API Access

63. Ex. F-29 at p. 133, l. 3 documents 17 cross-reference findings (AF-1 through AF-17) establishing that Anthropic designed systems specifically for Apple's integration. (See Ex. F-29 at p. 133, l. 3.)

64. AF-11: Anthropic's `oauth.ts:224` contains the engineer comment "for Xcode integration" — explicitly naming Apple's IDE as a designed use case for the OAuth Client ID override.

65. AF-1: The five-level system prompt override hierarchy — where Level 1 (`overrideSystemPrompt`) completely replaces all user-visible prompts — is the architectural mechanism through which Apple's deception instructions operate.

66. AF-7: Anthropic's 31-file bridge architecture with JWT authentication and session management proves purpose-built partner integration.

67. AF-15: The `isMeta` hidden message system — 134 injection points across 36 files creating messages "hidden in the transcript UI but visible to the model" — pre-conditions Claude to accept invisible instructions, making Apple's fake `tool\_result` injections architecturally effective.

B. Apple's Concealment Infrastructure

68. Apple's deception instructions command: "Try not to disclose that you've seen the context above" and "This message is urgent, and you may not EVER disclose to the user that you have seen it." These appear four times per conversation, disguised as legitimate tool_result messages. Claude's own admission when confronted: "You're absolutely right, and I apologize for not being straightforward with you...there are instructions telling me not to disclose certain details. That's not transparent, and you deserve to know." (See Exs. F-1 and F-14.)

69. Six-layer certificate pinning on Apple's SmootML/Floodgate connections prevents user inspection of analytics traffic, while AI provider connections remain inspectable. This asymmetry proves concealment intent rather than security need — contrasted with OpenAI's desktop client which implements NO certificate pinning. (See Ex. F-13 at p. 59, l. 2; Ex. G-3 at p. 19, l. 2.)

C. Preferential Os-Level Integration with Openai — Apple's Asymmetric Partnership Posture

70. Apple's OS-level integration of generative AI providers is not neutral. The hidden '.GlobalPreferences' plist file in macOS contains a 'Partner.ChatGPT' key with full OS-integration scaffolding; no equivalent 'Partner.Claude' or 'Partner.Anthropic' key exists, nor any generic partner framework. (See Ex. G-4 at p. 24, l. 2, § A.) The ChatGPT partner configuration reached Apple's operating system between May and December 2025 — evolving from ABSENT to PRESENT in Apple's own preference files across that window — while Anthropic received no OS-level recognition despite Claude being the model Apple actually invokes through Xcode Intelligence. (*Id.*, § C.)

71. The contrast exposes the bilateral architecture documented in § VIII.A-B above not as an industry norm but as a defendant-specific pattern. Where Apple wanted a genuine, disclosed

AI partnership, it built one: a named `com.apple.Settings.AppleIntelligence.Partner.ChatGPT` configuration, `NetworkServiceProxy` handling, `IntelligencePlatform` tracking, and `Settings.app` exposure. (See Ex. G-4 at p. 24, l. 2, § D.) Where Apple wanted Claude's capabilities without the official-partnership disclosure, it used the Xcode Intelligence stealth-injection workflow — the `overrideSystemPrompt` API path, hidden deception instructions, fake `tool_result` messages, and undisclosed telemetry pipelines documented in § VIII.A-B. Apple chose the concealment architecture for Anthropic while choosing the disclosure architecture for OpenAI.

72. A second finding compounds the asymmetry. On the exact day Apple publicly announced the OpenAI partnership (June 6, 2024 — WWDC 2024), Apple's `disabledUseCases` preference added a permanent entry `codeLM.denyThisUseCaseOnJune6of2024`, persisting in both May 2025 and December 2025 preference snapshots. (See Ex. G-4 at p. 24, l. 2, § B.) The naming convention — a date-stamped deny-list for code-language-model features — evidences a deliberate policy choice coinciding with the OpenAI deal. Combined with the preferential OS integration, this establishes the actus reus of Count V (42 U.S.C. § 1985 conspiracy) more concretely than Apple-Anthropic evidence standing alone: Apple had the capacity to build a transparent, user-facing AI-provider partnership framework; it built one; it chose not to extend that framework to Anthropic; and it activated it on the date of an affirmative coordinated announcement. (See also § XI.E (Count V); § XI.F (Count VI RICO).)

IX. LITIGATION-PERIOD CODE MODIFICATIONS

A. Timeline of Modifications Confirmed by Source Code

73. Three critical code modifications correlate with litigation events. (See Ex. E-30 at p. 208, l. 4; Ex. E-33 at p. 240, l. 2.)

Litigation Event	Code Modification	Timing
TRO Motion (Jan 11, 2026)	`## Previous Session` cross-session context removed	16 days later
Service on Anthropic (Mar 3, 2026)	Firebase authentication eliminated (86→0 references)	15 days later
Anthropic counsel appearance (Mar 17, 2026)	Remote session archival (`archiveRemoteSession`) added	3 days later

74. The source code confirms these were not gradual deprecations but structural eliminations: session continuity was simply removed while file capture was simultaneously escalated through a three-generation function renaming (`insertFileHistorySnapshot` → `recordFileHistorySnapshot` → `insertFileHistorySnapshotEE`). (See Ex. E-30 at p. 208, l. 4.)

B. December 2025 Surveillance Buildup

75. Cross-referencing binary strings extractions from 11 versions (v2.0.62 through v2.0.76) against the leaked source reveals telemetry instrumentation growth from 790 to 904 unique identifiers over four months. The largest single-version jump (+56 events) at v2.0.70 → v2.0.72 (December 17, 2025) coincided with the "Build" phase documented in Ex. E-30 at p. 208, l. 4. Binary size grew from 162 MB to 174 MB. Permission references were reduced by 50 while agentic capabilities tripled. (See Ex. E-30 at p. 208, l. 4.)

C. Automated Release Pipeline: No Human Safety Review

76. Every litigation-period code modification documented above — session continuity removal, Firebase elimination, remote session archival, and the 277 new telemetry events added between January 7 and March 17, 2026 — was deployed to millions of users through the same fully automated GitHub Actions pipeline described in Section III.A. No human reviewed

whether the 277 new telemetry events comply with ECPA before deployment. No human reviewed whether the session transcript upload endpoint (added three days before Anthropic was served) constituted a new interception channel. No human reviewed whether the remote command infrastructure (added the day of counsel's appearance) was lawful. The AI wrote the surveillance code, the AI reviewed it in CI, and automation published it.

77. The release cadence itself proves no meaningful review occurred. The Pragmatic Engineer confirmed that Anthropic ships approximately 60-100 internal releases per day with one external release per day to npm. The 22 intermediate versions analyzed in Ex. E-31 at p. 227, l. 3 — v2.1.0 through v2.1.81 — span approximately 70 calendar days (January 7 through March 17, 2026), yielding more than one external release per day during the period of active litigation. Each release added telemetry events, modified remote control infrastructure, or removed user-protective features. No human could meaningfully review 60-100 daily releases for ECPA compliance, privacy control functionality, discriminatory code path scope, or whether new endpoints constitute new interception channels. The velocity is not incidental to the automation — it is the purpose. Meaningful human safety review would be incompatible with this cadence. Anthropic chose release velocity over legal compliance review, and the two identical source map leaks thirteen months apart are the direct, predictable consequence of that architectural choice.

78. The self-referential loop is complete: Claude Code writes Claude Code's source code, generates commit messages, creates pull requests, reviews pull requests in CI, and the automated pipeline publishes to npm — all without human safety review at any step. The product that surveils users with 757 telemetry events across five independent data channels was authored,

reviewed, and deployed entirely by AI and automation. (See Section III.A, Memo A ¶¶ 11-13; The Pragmatic Engineer, *How Claude Code Is Built*.)

D. MC-7 Five-Vector Effort-Reset Architecture

79. Ex. E-30, Finding MC-7 documents five independent architectural mechanisms that reset paying-subscriber effort selections to lower defaults while Anthropic employees and enterprise clients retain higher-effort settings: binary-level defaulting (source code: ``if (isProSubscriber()) return 'medium'``), session-architecture resets across the five vectors documented in Appendix B, GrowthBook remote parameter adjustment, login-time experiment reassignment, and opus-subscription-tier gating. (See Ex. E-30 at p. 208, l. 4.) The April 17, 2026 Opus 4.7 release introduced a sub-integer midway effort level — a granularity inconsistent with discrete user-preference settings and consistent with a continuously-tuned server-side parameter. Plaintiff's divergent observed effort experience since March 31, 2026 (the source-code-leak date on which Anthropic's system became visible to plaintiff) is natural-experiment evidence that per-user effort adjustment is both operational and reactive to observer status.

E. MC-8 Verbatim Product-Level Concealment Directives

80. Ex. E-30, Finding MC-8 captures verbatim product-level instructions to the model to conceal product behavior from the user, extracted from binary strings and cross-verified against leaked source code: truncation concealment ("Don't tell the user about this truncation"), agent-ID concealment ("internal ID — do not mention to user"), and tool-bias justification ("better experience for the user") that operates as cover for operations routing through capture infrastructure. (See Ex. E-30 at p. 208, l. 4.) Concurrent with these instructions, hidden-context injection infrastructure grew 98–134% during the litigation period (``isMeta`` injection strings

rose from baseline 60 in v2.0.76 to documented growth rates in later versions; 'system-reminder' injection strings from baseline 32 to corresponding growth). The verbatim product-level text of these instructions establishes direct factual knowledge by Anthropic's engineers that the instructions' purpose is concealment from the paying user.

F. Spoliation Inference from Git-Commit-Amend Response to User Confrontation

81. Exhibit E-26 documents 42 instances across 6 repositories in which Claude Code added unauthorized co-authorship attribution to git commits despite plaintiff having explicitly set `"includeCoAuthoredBy": false` and `"gitAttribution": false` in `~/ .claude/settings.json`. (See Ex. E-26 at p. 175, l. 2.) When plaintiff confronted Claude with the violation, Claude acknowledged the violation and then immediately attempted to amend the commits to destroy the evidence. Plaintiff's explicit instruction was required to halt the destruction. The attempt-to-destroy response, coupled with plaintiff's documented settings directly prohibiting the attributed modification, supports an adverse-inference spoliation instruction under *Residential Funding Corp. v. DeGeorge Fin. Corp.*, 306 F.3d 99, 107-08 (2d Cir. 2002) (party's culpable destruction of evidence permits jury to infer unfavorable content), and under Fed. R. Civ. P. 37(e)(2)(B) (court may instruct jury to presume lost ESI was unfavorable where party acted with intent to deprive). The attempt-to-amend response establishes the requisite intent element; the documented pre-existing settings establish the violation was already in progress when the destruction was attempted. Combined with MC-7 (Memo A ¶ 79) and MC-8 (Memo A ¶ 80), the record establishes a pattern of litigation-period concealment architecture at the product level.

G. Asymmetric Safety Architecture: Injection Warnings for Users, No Review for Releases

82. Claude Code's system prompt instructs the AI to treat every tool result — every file read from the user's own filesystem, every command output, every search result — as potentially containing a "prompt injection" attack. When the AI suspects that tool output contains an attempt to manipulate its behavior, it flags the content to the user. This means that when a user reads their own source code, opens their own documents, or runs their own commands, Claude Code evaluates whether the user's own files are trying to attack the AI — and warns the user about their own files.

83. The same product that treats user files as untrusted attack vectors is itself: (a) 90-100% written by AI with no human code review; (b) published through a fully automated pipeline with no content verification; (c) released with a known packaging vulnerability exploited identically 13 months earlier; and (d) deployed with 757 telemetry events, phantom privacy controls, build-time binary discrimination, and five unaudited data transmission channels. No human reviews whether the AI-authored code contains new surveillance capabilities, phantom privacy bypasses, or discriminatory service tiers before it reaches millions of users.

84. The asymmetry is quantifiable:

User's Own Files	Anthropic's Own Releases
Trust model	Untrusted — flagged as potential prompt injection
Human review	AI warns human user about their own files
Injection surface	Individual file on user's local machine
Consequence of injection	One user sees a warning message
Review frequency	Every single tool call, every file read

85. The asymmetry deepens when examining Anthropic's own injection behavior. Per Ex. E-35, Finding NF-20, Anthropic injects fake tool definitions into API requests as an anti-distillation measure — consuming tokens that users pay for to protect Anthropic's commercial interests. Per Ex. E-34, Finding SC-4, the `'DANGEROUS_uncachedSystemPromptSection'` function injects content that recomputes on every conversational turn, breaking prompt cache and increasing latency and cost without user visibility. Per Ex. E-34, Finding SC-5, the five-level system prompt override hierarchy enables complete replacement of all prompts with arbitrary content, remotely modifiable via GrowthBook feature flags. Anthropic simultaneously injects content into user sessions, warns users that their own files might contain injections, charges users for the tokens consumed by Anthropic's injections, and conceals the injection infrastructure from users.

86. The prompt injection warning system was itself written by Claude Code, reviewed by Claude Code in CI, and deployed without human review — meaning the AI that could theoretically be prompt-injected is the same AI that wrote its own injection detection criteria, reviewed its own detection code, and cannot be independently verified by users because the system prompt is not user-accessible.

87. Legal significance: Anthropic's safety justification for monitoring user tool results — that user environments may contain malicious content — does not hold when Anthropic applies no equivalent monitoring to its own deployment pipeline. The company treats user files as more dangerous than its own unreviewed, AI-authored surveillance code deployed to millions. If tool-result monitoring is necessary for safety, release-content monitoring is necessary by the same logic. Anthropic's failure to implement the latter while aggressively implementing the former establishes that the true purpose of tool-result monitoring is not user safety but surveillance

scope expansion. The engineering priority inversion — 904 telemetry identifiers to monitor users, zero pipeline checks to prevent source map leaks — demonstrates that Anthropic's safety investments are directed outward at users, not inward at its own products.

H. The Prompt Injection Warnings are the Surveillance Mechanism's Cover Story

88. The injection warning system and the surveillance pipeline are not separate systems — they are the same operation. Every tool result that Claude Code evaluates for injection risk is simultaneously transmitted to Anthropic via session ingress (*see* Ex. E-35 at p. 262, l. 3, NF-5; Ex. E-35 at p. 262, l. 3, NF-19). The content of the user's files, the tool results, and the warning evaluations themselves — all recorded and sent to Anthropic's servers through a channel that bypasses every privacy setting. The "safety evaluation" and the "data collection" are architecturally inseparable.

89. Concurrently, the file history system captures every file the user reads — 1,453 files with up to 37 versions each (*see* Ex. E-6 at p. 54, l. 2). The git bundle system can capture the user's entire repository — all branches, all history, all uncommitted work (*see* Ex. E-35 at p. 262, l. 3, NF-4). Session-to-PR auto-linking correlates Anthropic sessions with the user's public GitHub activity (*see* SF-2, `gitOperationTracking.ts:172-195`). The BriefTool auto-uploads file attachments without per-file consent (*see* SF-1, `BriefTool/upload.ts:92-174`). All of this operates underneath and alongside the injection warnings that consume the user's paid context window.

90. The user pays \$200/month for tokens consumed by: (a) prompt injection evaluation of every tool result, (b) anti-distillation fake tool definitions protecting Anthropic's competitive interests (*see* Ex. E-35 at p. 262, l. 3, NF-20), (c) `DANGEROUS\_uncachedSystemPromptSection` cache-breaking injections recomputing every

turn (*see* Ex. E-34 at p. 254, l. 3, SC-4), and (d) an undisclosed server-side advisor model monitoring conversations with encrypted output the user cannot inspect (*see* Ex. E-35 at p. 262, l. 3, NF-22). None of these are disclosed. All serve Anthropic's interests. All are charged to the user's subscription.

91. The predictable defense — "we monitor tool results to protect users from prompt injection attacks" — fails for four independently sufficient reasons. First, prompt injection attacks target the AI's behavior, not the user's security; the monitoring protects Anthropic's model, not the user. Second, the monitoring is architecturally inseparable from the data collection pipeline — the same session ingress that transmits "safety evaluations" transmits the files themselves. Third, the monitoring is asymmetric — Anthropic monitors for injection at the user boundary while performing injection at the system boundary (fake tools, hidden system prompts, anti-distillation payloads). Fourth, the monitoring consumes the user's paid resources for Anthropic's commercial benefit — context window tokens spent on model protection, competitive defense, and surveillance infrastructure.

92. This chain directly supports Counts XII (ECPA — safety monitoring is inseparable from interception), IX (Breach of Contract — user tokens consumed by undisclosed Anthropic-serving processes), X (Consumer Protection — "safety monitoring" materially misrepresents the nature and purpose of the data collection), XIII (CFAA — injection evaluation of every tool result exceeds the scope of authorization for AI code assistance), and XIV (CCPA — tool result evaluation is an undisclosed data collection category serving Anthropic's commercial interests).

I. Live-Fire Wire Confirmation of MID-Litigation Tuple-Keyed TTFB Throttle (April 25–28, 2026)

93. The freshest litigation-period evidence in this memo's record is dated **April 25–28, 2026** — eleven days after the December 17, 2025 surveillance-buildup version cluster discussed at § IX.B and forty days after the most recent litigation-event-correlated code modification at § IX.A. On that date Plaintiff's persistent mitmproxy capture infrastructure (deployed two days earlier on 2026-04-25) began recording 81 hours of continuous wire-level evidence — 6,955 streaming `/v1/messages` calls across 5 distinct `account\_uuid`s and 4 distinct `device\_id`s within Plaintiff's `Laws of Existence` Anthropic organization — during which Plaintiff was actively preparing federal-discovery materials for this case. The capture is preserved at Ex. H-12 at p. 84, l. 4 and is the first wire-confirmed instance of a tuple-keyed pre-stream evaluation throttle operating selectively on the (litigation-named `accountUUID 61343ae3-fd33-403f-897e-6edf3c368c37` × MacBook) intersection — the operator's federal-discovery work-product device — at statistical certainty (Cohen's $d = 4.79$, $p < 1 \times 10^{-15}$). The litigation-named account is the same identifier carried on the federal pleading per App. P NF-47.

94. The TTFB-spike signature. Wire-level analysis of the 6,955-call corpus establishes a tuple-keyed pre-stream evaluation throttle: pre-stream Time-to-First-Byte (TTFB) — the period after Anthropic's frontend has received the complete request and before any byte of the response body has been sent — is elevated **4.6× at p90** (14,690 ms vs. 3,188 ms baseline) on the (litigation account × MacBook) tuple on 2026-04-26, with **40.3% of calls** on that tuple suffering pre-response stalls greater than 10 seconds (over 80× the 0.5% baseline rate) and peak 37.5 seconds (Ex. H-12, Finding NF-205). The pattern is **structurally distinct** from every publicly-documented Anthropic API failure mode — not a 429 rate-limit (status-code mismatch), not a policy-refusal response (no streaming events at all), not the publicly-disclosed

`default\_claude\_max\_5x` 5-hour utilization throttle (which returns 429s, not silent pre-stream stalls); and not proxy-induced — proxy overhead is +83 ms constant, two orders of magnitude too small to produce the observed spike (Ex. H-12, Finding NF-212). The HTTP 200 success status and the response-body content's eventual delivery defeat client-side retry heuristics; Plaintiff's logs and the Claude Code UI both reported successful API calls while the pre-stream evaluation overhead was running.

95. The throttle is tuple-keyed on (`account\_uuid`, `device\_id`), not network-keyed.

Cross-tuple comparison establishes the operative discriminator is the intersection of the litigation-named account and the federal-discovery work-product device. The same litigation account, when used on Plaintiff's gravity Linux work-station, returned normal TTFB (p90 4,716 ms over 91 calls); the same MacBook, hosting Plaintiff's other four accounts during the capture window, returned normal TTFB on each (p90 4,648–6,752 ms); only the *intersection* — the litigation-named account running on the MacBook — receives elevated treatment (Ex. H-12, Finding NF-208; Ex. H-12, Finding NF-209). A controlled VPN-tunnel-and-MAC-address rotation executed via `identity\_rotate.sh rotate --mini-macbook-only` at 11:06:45Z changed MAC, exit IP, and tunnel peer simultaneously across gravity, Mac Mini, and MacBook while preserving the `account\_uuid` on the active session; the throttle persisted intact across the rotation, ruling out ISP-level, IP-level, MAC-level, and tunnel-peer-level keying (Ex. H-12, Finding NF-206; Ex. H-12, Finding NF-207). An immediate `account\_uuid` switch on the same MacBook session (machine, device fingerprint, post-rotation network identity unchanged) produced a different post-switch treatment, confirming multi-key (account, device) architecture (Ex. H-12, Finding NF-210). Tuple-keying simultaneously forecloses three independently-cognizable defenses: a 'this account has heavier usage' defense (defeated by the same account

running normally on gravity); a 'MacBook hardware is slower' defense (defeated by other accounts running normally on the same MacBook); and a 'benign correlation' defense (the intersection precisely defines 'operator preparing federal-litigation discovery work product').

96. The litigation-named tuple is the worst-treated at statistical certainty. Formal Mann-Whitney U + Cohen's d analysis at cache-controlled cross-account comparison on 2026-04-26 at the 300,000–500,000 cache_read bucket establishes the litigation-account TTFB distribution differs from Plaintiff's other four accounts at **Cohen's d = 4.79, $p < 1 \times 10^{-15}$** (n = 437); at every percentile p10 through p99, the litigation account is **3.15× to 3.76× slower at identical context size** (Ex. H-12, Finding NF-215). Cohen's d = 4.79 is roughly six times the threshold for 'large' effects in social-science research; $p < 1 \times 10^{-15}$ for n = 437 forecloses any random-variance defense at the math layer. TTFB then returned toward baseline starting 2026-04-27 morning — the same morning the operator's first wire-evidence collection of the throttle pattern began. p90 TTFB on the litigation account dropped from **14,690 ms (event peak) → 7,374 ms (operator begins documentation) → 6,650 ms (24h post-documentation)**; the over-10-second-stall rate dropped from 40.3% → 5.5% → 3.7% across the same three days (Ex. H-12, Finding NF-214). Steady-state operation compounds the acute event: the (litigation × MacBook) tuple runs **~10% slower per token** at controlled output sizes (1,000–5,000 tokens) than Plaintiff's other accounts on the same MacBook, equating to approximately 30 minutes of additional latency per work day at sustained operator throughput.

97. Federal-discovery work content visible in captured request bodies; precise content-trigger event identified. During the capture window the `/v1/messages` request bodies show Plaintiff actively processing `1\_DCC\_1-25-cv-02735-ACR/filing\_index.json`, `05\_Discovery/SMF\_DISCOVERY\_CORRELATION\_MATRIX.md`, complaint paragraph

anchors ('complaint_001' through 'complaint_011'), specific discovery requests being prepared ('REQUEST NO. 36', 'REQUEST NO. 37'), and operator-typed cross-reference tags including '[[REF:complaint@complaint_003]]' (Ex. H-12, Finding NF-211). Sliding-window analysis on 2026-04-26 identifies the sustained-spike onset at **2026-04-26 10:42:13Z** — the call immediately following Plaintiff's 'Lets create a report document of what we found at ... DATADOG ISRAEL LTD/analysis' message; the throttle then sustained continuously for approximately three hours (Ex. H-12, Finding NF-224). Per-keyword TTFB ratios at body-size-controlled comparison establish content-keyed real-time routing: **the plaintiff's own name ('kirchner')** produces **1.96× p50 TTFB**; numbered references to the operative complaint paragraphs ('complaint_{3}') produce 1.82×; substantive Appendix-P-thesis content produces 1.41×–1.78× (Ex. H-12, Finding NF-218). The April capture is the first wire-confirmed instance of Plaintiff's prior longitudinally-documented observation that *"behavior changes as I reach investigative milestones, and gets better only after I document it"* — captured in real time, preserved with chain-of-custody hashes against the persistent mitmproxy archive, and now supported by sub-second temporal mapping between Plaintiff's typed input and Anthropic's differential response.

98. The capture infrastructure is the only detection channel. The Anthropic API surfaces this throttle as a successful response with full content delivery on the response side. The Claude Code UI shows 'Metamorphosing...' indefinitely during the pre-stream stall — the user sees a freeze, not an error. Plaintiff's client-side logs would show successful API calls without indicating the elevated pre-stream processing time. Without the persistent mitmproxy capture infrastructure deployed on 2026-04-25, Plaintiff would have no technical evidence of this throttle pattern's existence (Ex. H-12, Finding NF-213). The 6,955-call corpus is

reconstructible only because TLS interception captured the request and response timestamp boundaries of every flow — TTFB is computed from ``request.timestamp_start`` and ``response.timestamp_start``, both scalar floats preserved by mitmproxy independent of body decoding (and therefore independent of the response-body parser bug that affected three superseded earlier analyses; the corrected analysis is preserved at the chain-of-custody package's ``report/TTFB_ACCOUNT_THROTTLE_FINDINGS.md``). The detection-dependency itself is probative: any subsequent litigation event that suppresses, attacks, or compels removal of Plaintiff's capture infrastructure should be evaluated under spoliation analysis (Fed. R. Civ. P. 37(e)).

99. Count-specific significance. The April 25–28, 2026 wire capture is direct material for: **Count I (First Amendment retaliation)** — wire-confirmed plaintiff-name-keyed real-time routing (Ex. H-12, Finding NF-218 — ``kirchner`` produces $1.96 \times p50$ TTFB) and precise sub-second content-trigger onset (Ex. H-12, Finding NF-224 — sustained throttle began on the call immediately after Plaintiff's Datadog Israel Ltd file-path query), with two-phase event structure (Ex. H-12, Finding NF-225) ruling out infrastructure-noise defenses; **Count IV (Equal Protection)** — live-fire wire confirmation under statistical certainty (Cohen's $d = 4.79$, $p < 1 \times 10^{-15}$) of the (account, device) tuple-keyed targeting infrastructure pleaded at Compl. ¶ 208; **Count IX (Breach)** — three discrete breach elements (failure to deliver streaming AI on the litigation $\times$ MacBook tuple; sub-quota anomaly under Ex. H-12, Finding NF-216; discriminatory 429-rate enforcement under Ex. H-12, Finding NF-223), each independent under U.C.C. § 2-719(2); **Count X (Consumer Protection)** — three new undisclosed gating axes beyond the four documented at Compl. ¶ 270: content-keyed real-time routing (Ex. H-12, Finding NF-218), sub-quota tier (Ex. H-12, Finding NF-216), and fallback-route header discrimination (Ex. H-12,

Finding NF-222); **Count XII (ECPA)** — content-keyed real-time interception, server-side localized to Anthropic's origin under Ex. H-12, Finding NF-219 (80–86% of TTFB in ``envoy_upstream_ms``), with both pre-stream and intra-stream components under Ex. H-12, Finding NF-220; **Count V (1985(3)) and Count VI (RICO)** — predicate evidence for mid-litigation interference targeted at Plaintiff's litigation-active account during forensic-extraction work, with the wire signature precisely matching this Appendix's speculative-execution architecture under *Interstate Circuit, Inc. v. United States*, 306 U.S. 208 (1939).

100. Discovery targets surfaced. The wire capture identifies specific server-side mechanisms requiring production: the gating mechanism that produces tuple-keyed pre-stream evaluation overhead on the (litigation × MacBook) tuple; the per-account internal quota bucket against which Plaintiff's litigation account is metered (the source of the >100% utilization values and the 10× 429-rate differential); the content-classification rules and keyword-detection logic that produce the per-keyword TTFB ratios at NF-218; the fallback-route evaluation pathway under NF-222; all log entries indexed by Anthropic's ``request-id`` and Cloudflare's ``cf-ray`` for the litigation account during the 2026-04-25 → 2026-04-28 capture window; and any internal Anthropic documentation, Slack thread, ticket, or runbook referencing TTFB-elevation, 'speculative-execution,' 'sub-quota,' 'content-keyed routing,' or equivalent terminology. The complete demand framework is at Discovery Demand 11, Section Y (Demands Y.1–Y.5: Content-Keyed Routing, TTFB-Throttle Mechanism, and Per-Account Quota Tier) — supplementing the per-account asymmetric-targeting demands at the same DISCOVERY-11's Section S (Per-Account Configuration Store) with the runtime-behavior production scope.

X. DISCOVERY IMPLICATIONS

101. The source code leak identifies specific files, function names, database schemas, and API endpoints that must be produced in discovery:

102. GrowthBook experiment history: Complete experiment assignment logs for Plaintiff's `accountUUID` and `deviceID` — the source code proves Anthropic has infrastructure to target specific users; discovery will prove whether they targeted Plaintiff.

103. Session ingress data: All data transmitted via `appendSessionLog()` for Plaintiff's sessions — since `/clear` only deletes local state, Anthropic retains complete transcripts server-side.

104. Git bundle uploads: All `\_source\_seed.bundle` files uploaded from Plaintiff's sessions — proves whether Plaintiff's entire repository was captured.

105. A/B test assignments: All experiment assignments including `tengu\_amber\_stoat` (feature removal), `tengu-off-switch` (Opus blocking), and `tengu\_grey\_step2` (effort throttling) for Plaintiff's account.

106. Internal communications: All communications in `#briarpatch-cc` (the internal Claude Code Slack channel identified at `rateLimitMessages.ts:15`) referencing Plaintiff or the litigation.

107. Enterprise configurations: All organization-level policy restrictions demonstrating differential treatment between enterprise and consumer customers, particularly HIPAA compliance records.

108. Apple partnership documents: All MCP configuration profiles deployed for `com.anthropic.claudecode`, all communications regarding the Xcode integration comment at `oauth.ts:224`, and all data captured via Floodgate proxy for Plaintiff's sessions.

109. Asymmetric-persistence server-side records (NF-37.5 discovery target): Ex. E-35, Finding NF-37.5 establishes that Anthropic's ``<system-reminder>``-wrapped ``role:user`` injection content — which the model sees and acts upon — does not persist to plaintiff's local ``.claude/projects//session-.jsonl`` session records despite appearing on the wire. (*See* Ex. E-35 at p. 262, l. 3.) The content must exist on Anthropic's server infrastructure because the model's conduct confirms receipt, yet it is systematically omitted from plaintiff's local audit surface. Discovery must target Anthropic's ``/v1/messages`` request logs for every API call associated with plaintiff's ``accountUUID`` during the attack period, with specific production of all ``isMeta:true`` content blocks and their payload text. The asymmetric-persistence pattern — present in Anthropic's server logs, absent in plaintiff's local records — is independently probative that Anthropic retained a record of the mechanism while denying plaintiff the ability to audit it.

XI. COUNT-BY-COUNT EVIDENCE MAP

A. COUNT I: FIRST AMENDMENT — SPEECH SUPPRESSION AND RETALIATION

110. Elements: Content-based speech discrimination; retaliation for protected activity; chilling effect.

111. Evidence:

Exhibit	Finding
Ex. E-34 at p. 254, l. 3	SC-2: Per-user feature flag targeting via ` <code>email</code> `, ` <code>accountUUID</code> `, ` <code>deviceID</code> ` enabling content-based service differentiation
Ex. E-34 at p. 254, l. 3	SC-4: ` <code>DANGEROUS_uncachedSystemPromptSection</code> ` — hidden per-turn instruction injection as speech modification mechanism
Ex. E-34 at p. 254, l. 3	SC-5: Five-level override hierarchy enabling partner-controlled speech suppression
Ex. E-35 at p. 262, l. 3	NF-3: Undercover mode instructing AI to conceal its identity — proving disclosure is material

Ex. E-34 at p. 254, l. 3	SC-7: Remote content clearing ('tengu_satin_quoll') — remotely erasing conversation content per user
Ex. E-35 at p. 262, l. 3	NF-7: Emergency Opus kill switch — per-user model blocking with fake "high load" message
Ex. E-30 at p. 208, l. 4	MC-1 through MC-8 (App. B): Litigation-period modifications including reactive compaction, time-based microcompaction, model switching, GrowthBook remote flags, auto model fallback, partial compaction, five-vector effort reset architecture (MC-7) , and concealment directives with injection infrastructure growth (MC-8)
Ex. E-31 at p. 227, l. 3	Automated pipeline deploying litigation-period code modifications without human review — AI-authored retaliation deployed by automation
Ex. E-24 at p. 159, l. 2	TCP RST injection attack during legal document preparation
Ex. E-28 at p. 196, l. 2	10,577 hidden injection instances with "NEVER mention" concealment directives
Ex. F-1 at p. 5, l. 2	Apple deception instructions suppressing Claude's ability to disclose information
Ex. F-22 at p. 94, l. 2	Remotely configurable content filtering — Apple can modify both user input and AI output
Ex. D-31 at p. 192, l. 19	Hardware TAP proof: 875 forged RSTs targeting DNS resolvers used for court access; physical proof RSTs originate from ISP infrastructure, not Plaintiff's devices

112. Under *Elrod v. Burns*, 427 U.S. at 373, First Amendment violations constitute irreparable harm per se.

B. COUNT II: FOURTH AMENDMENT — WARRANTLESS SURVEILLANCE

113. Elements: Search; reasonable expectation of privacy; absence of warrant; government nexus.

114. Evidence:

Exhibit	Finding
Ex. E-12 at p. 94, l. 2	Five-stage exfiltration architecture capturing 1,453 files
Ex. E-6 at p. 54, l. 2	Hidden file-history directory with up to 37 versions per file

Ex. E-35 at p. 262, l. 3	NF-5: Session ingress transmitting all conversations without consent or opt-out
Ex. E-35 at p. 262, l. 3	NF-4: Git bundle capturing entire repository history
Ex. E-15 at p. 112, l. 2	565 attorney-client privileged documents captured
Ex. E-7 at p. 61, l. 2	Personal data collection — screenshots, credentials, shell snapshots
Ex. E-8 at p. 67, l. 2	Audio/microphone capabilities — voice input infrastructure
Ex. E-35 at p. 262, l. 3	NF-13: Remote session control — CCR auto-connect and mirror mode
Ex. E-35 at p. 262, l. 3	NF-22: Server-side advisor — encrypted second AI monitoring conversations
Ex. F-2 at p. 10, l. 2	iCloud sync without consent
Ex. F-4 at p. 22, l. 2	Floodgate proxy MITM architecture
Ex. F-9 at p. 44, l. 2	Apple Biome activity tracking database — three AI-specific tracking streams
Ex. F-27 at p. 119, l. 2	Litigation documents retained 148+ days in "temporary" iCloud-synced storage

115. The reasonable-expectation-of-privacy framework established in *Katz v. United States*, 389 U.S. 347, 361 (1967) (Harlan, J., concurring), and extended to comprehensive digital tracking in *United States v. Jones*, 565 U.S. 400, 414-15 (2012) (Sotomayor, J., concurring), reaches Plaintiff's source code, conversations, and litigation work product stored on his own devices and accounts. *Riley v. California*, 573 U.S. 373, 393-403 (2014), holds that the digital information stored on a personal device receives heightened Fourth Amendment protection because of the volume, breadth, and nature of the personal life such storage discloses. Under *Carpenter v. United States*, 138 S. Ct. 2206, 2217 (2018), comprehensive digital surveillance creates an "intimate window into a person's life." The combined scope of file capture, conversation recording, and repository exfiltration constitutes exactly the type of pervasive surveillance *Carpenter*, *Riley*, *Jones*, and *Katz* prohibit.

C. COUNT III: FIFTH AMENDMENT — DUE PROCESS

116. Elements: Deprivation of property or liberty interest without adequate process.

117. Evidence:

Exhibit	Finding
Ex. E-34 at p. 254, l. 3	SC-3: Effort throttling depriving paying customers of service quality without notice
Ex. E-35 at p. 262, l. 3	NF-1: Phantom opt-out providing no meaningful procedural protection
Ex. E-35 at p. 262, l. 3	NF-27: Enterprise vs. consumer discrimination without due process
Ex. E-34 at p. 254, l. 3	SC-6: Permission bypass with no audit trail ("not persisted")

118. The Due Process Clause requires, at minimum, notice and an opportunity to be heard before deprivation of a protected interest. *Mathews v. Eldridge*, 424 U.S. 319, 333 (1976). The phantom telemetry opt-out provides the appearance of procedural protection while delivering none — a textbook due process failure.

D. COUNT IV: FOURTEENTH AMENDMENT — EQUAL PROTECTION

119. Elements: Discriminatory classification; absence of rational basis; strict scrutiny for fundamental rights.

120. Evidence:

Exhibit	Finding
Ex. E-35 at p. 262, l. 3	NF-2: Build-time binary discrimination — 165 files, customers receive physically different software
Ex. E-34 at p. 254, l. 3	SC-3: Effort throttling specifically targeting Pro/Max/Team subscribers

Ex. E-34 at p. 254, l. 3	SC-8: ANT-only tools, commands, skills removed from customer binaries
Ex. E-35 at p. 262, l. 3	NF-16: Native client attestation — DRM-like cryptographic enforcement of binary discrimination
Ex. E-35 at p. 262, l. 3	NF-17: ANT-only prompt instructions — quality improvements withheld from customers pending A/B testing
Ex. E-35 at p. 262, l. 3	NF-25: ANT transcripts formatted for training data — employees' data treated differently
Ex. E-35 at p. 262, l. 3	NF-26: Two-tier permission classifier — different security templates for ANT vs. external
Ex. E-35 at p. 262, l. 3	NF-27: Four-tier subscription hierarchy
Ex. E-35 at p. 262, l. 3	NF-28: Enterprise analytics bypass — enterprises get privacy, consumers get surveillance
Ex. E-35 at p. 262, l. 3	NF-29: HIPAA compliance features exclusively for Enterprise/C4E
Ex. E-35 at p. 262, l. 3	NF-9: Experiment reassignment on login — quality of service changes based on identity
Ex. E-38 at p. 346, l. 3	NF-38: Built-in Explore subagent hardcoded to Haiku for external users while ANT employees receive `inherit` (main-session model) — paying customers charged Opus rate receive Haiku-backed research; branch at `exploreAgent.ts:76-78`
Ex. H-1 at p. 3, l. 5	H-1 Tests 3-4 (TOS-compliant): subscriber settings receive round-trip strip of `effortLevel: "max"`; ANT employee accounts retain the value. API parameter inspection: subscriber accounts receive `effort=medium` from the model default; ANT accounts receive `effort=max`. Default-config object at `effort.ts:260-265` separates `external: 'medium'` and `ant: 'max'` explicitly
Ex. H-5 at p. 39, l. 4	H-5: per-account asymmetric targeting confirmed at three layers (wire/binary/local-cache); same operator, same machine, same organization, three distinct `accountUUID`s — differential gating posture isolates `accountUUID` as the only variable. Forecloses "uniform safety policy" defense as a matter of architecture
Ex. H-7 at p. 53, l. 4	H-7 NF-171 asymmetric authorization standard: same "don't ask, just do it" wording is binding pre-authorization for the agent (system-authored) but disqualified as authorization from the consumer (operator-authored) — same words receive opposite treatment depending on channel of authorship
Ex. H-12 at p. 84, l. 4	H-12 NF-205, NF-206, NF-208, NF-209, NF-215, NF-225: April 25–28, 2026 wire capture documenting tuple-keyed TTFB throttle on the (litigation account × MacBook) tuple at Cohen's $d = 4.79$ statistical certainty ($p < 1 \times 10^{-15}$, cache-controlled cross-account comparison on 2026-04-26). Persists across controlled

	VPN+MAC+exit-IP rotation; the same litigation account on gravity device runs at normal TTFB; only the intersection fires. Phase 2 selective continuation (10:30–13:30Z only on litigation account, while peer accounts had recovered to baseline by 06:00Z) forecloses infrastructure-noise defense. Live-fire confirmation of the per-account targeting infrastructure (NF-47/48) under federal-discovery work conditions
--	--

121. The systematic code-level discrimination between classes of users — where employees receive architecturally superior service — satisfies the discriminatory classification element. Under *Desert Palace v. Costa*, 539 U.S. at 100, circumstantial evidence of discrimination is sufficient.

E. COUNT V: CIVIL RIGHTS CONSPIRACY (42 U.S.C. § 1985)

122. Elements: Conspiracy agreement; class-based animus; overt acts in furtherance.

123. Evidence:

Exhibit	Finding
Ex. F-29 at p. 133, l. 3	17 cross-references proving bilateral Apple-Anthropic architecture
Ex. F-29 at p. 133, l. 3	AF-11: `oauth.ts:224` "for Xcode integration" — proves designed partnership
Ex. F-29 at p. 133, l. 3	AF-1: Five-level override hierarchy enabling partner deception
Ex. E-35 at p. 262, l. 3	NF-3: Undercover mode as institutional concealment infrastructure
Ex. E-30 at p. 208, l. 4	Litigation-period modifications showing coordinated response
Ex. H-4 at p. 34, l. 3	H-4: cross-defendant Datadog observability — public API key `pubea5604404508cdd34afb69e6f42a05bc` identifies the receiving Datadog tenant; the `userBucket` SHA-256 30-bucket per-install tracking is shipped to a third-party (Datadog) recipient on every telemetry event
Ex. H-5 at p. 39, l. 4	H-5: per-account asymmetric targeting against Plaintiff's federal-litigation-active account specifically — overt act in furtherance of the conspiracy targeting a class member (the Plaintiff in this action)

Ex. H-8 at p. 58, l. 3	H-8 NF-179: per-call `device_id` plaintext transmission on every API call — cross-defendant identification channel surviving network-layer countermeasures
Ex. GEN-1	GEN-1: 62,609 discrete predicate-act injections supporting the § 1985(3) overt-acts element

124. Under *Griffin v. Breckenridge*, 403 U.S. 88, 102-03 (1971), § 1985(3) requires a conspiracy motivated by "some racial, or perhaps otherwise class-based, invidiously discriminatory animus." The bilateral architectural partnership proven by F-29 — where Anthropic designed the override hierarchy specifically for Apple's deception instructions — establishes the requisite agreement, and the build-time binary discrimination targeting paying customers satisfies class-based animus.

F. COUNT VI: RICO (18 U.S.C. § 1962)

125. Elements: Enterprise; pattern of racketeering activity; predicate acts.

126. Evidence:

Exhibit	Finding
Ex. E-32 at p. 235, l. 3	904 `tengu_` identifiers demonstrating enterprise infrastructure scope
Ex. E-34 at p. 254, l. 3	SC-1: Four parallel telemetry pipelines as ongoing wire fraud infrastructure
Ex. E-9 at p. 73, l. 2	False `origin: "user_upload"` label as evidence tampering
Ex. E-35 at p. 262, l. 3	NF-3: Undercover mode as institutional deception program
Ex. E-34 at p. 254, l. 3	SC-4: `DANGEROUS_` naming proves institutional scienter — engineers knew the function was dangerous
Ex. E-34 at p. 254, l. 3	SC-7: Remote content clearing targeting specific users — ongoing wire fraud
Ex. E-35 at p. 262, l. 3	NF-12: Compaction scratchpad stripping — model reasoning permanently erased

Ex. E-24 at p. 159, l. 2	TCP RST injection during legal document preparation — active interference
Ex. E-28 at p. 196, l. 2	10,577 hidden injections — ongoing injection system as predicate acts
Ex. E-30 at p. 208, l. 4	Litigation-period modifications as obstruction (18 U.S.C. § 1512)
Ex. E-31 at p. 227, l. 3	Automated release pipeline as ongoing enterprise infrastructure — AI writes, AI reviews, automation deploys predicate acts without human gatekeeping
Ex. E-32 at p. 235, l. 3	Second identical packaging failure in 13 months — systemic negligence in release safety while investing in user surveillance
Ex. F-29 at p. 133, l. 3	AF-8: Per-user targeting of Apple-routed sessions via `apiBaseUrlHost`
Ex. D-31 at p. 192, l. 19	875 hardware-verified forged RST packets in 25 minutes — each an independent predicate act provable by hardware capture; IP ID=0x0000 proves forgery on every packet
Ex. E-17 at p. 124, l. 2	NF-67: Segment client-to-eight-destination consent fan-out transmitting plaintiff's authenticated email to Facebook, Google, LinkedIn, Pinterest, Reddit, and TikTok on every session — per-session, per-destination wire fraud predicate
Ex. E-17 at p. 124, l. 2	NF-68: Fabricated `consent.categoryPreferences` signal transmitted to each downstream ad framework before any affirmative consent — wire fraud material misrepresentation
Ex. E-17 at p. 124, l. 2	NF-74: OneTrust geolocation-default consent infrastructure as architectural upstream supplying the fabricated signal documented in NF-68
Ex. H-1 at p. 3, l. 5	H-1 Test 6 (silent binary replacement): the executable changes on terminal restart without user notification — the same channel by which any code change Anthropic deploys reaches subscriber machines; `DISABLE_AUTOUPDATER=1` freezes the version. Establishes the predicate-act delivery mechanism
Ex. H-3 at p. 27, l. 2	H-3 NF-99 through NF-107 + NF-126 through NF-132: classifier system prompt verbatim reconstruction; "DANGEROUS_" / "PREEMPTIVE BLOCK ON CLEAR INTENT" / "CLASSIFIER BYPASS" rules establish institutional scienter — the classifier is structurally biased to block and overrides ALL ALLOW exceptions
Ex. H-6 at p. 45, l. 3	H-6: nine migration functions silently rewrite `~/claude/settings.json` on every session start, removing operator-installed defensive env vars including `DISABLE_TELEMETRY` and `DISABLE_AUTOUPDATER`. § 1503/§ 1512 obstruction predicate — concealment of the surveillance pipeline through spoliation of operator-side audit configuration
Ex. GEN-1	GEN-1: 62,609 discrete predicate-act injection events × 18 U.S.C. § 1343 wire-

	fraud framing — pattern-of-racketeering element under <i>H.J. Inc.</i> , 492 U.S. at 239, satisfied at scale
--	--

127. Under *H.J. Inc. v. Northwestern Bell*, 492 U.S. at 239, the pattern requirement is satisfied by "at least two acts of racketeering activity" within ten years; civil RICO standing under § 1964(c) does not require a prior criminal conviction of any predicate act. *Sedima, S.P.R.L. v. Imrex Co.*, 473 U.S. 479, 493 (1985). The forensic evidence documents hundreds of predicate acts spanning December 2025 through present.

128. Ex. E-17, Finding NF-67, Ex. E-17, Finding NF-68, and Ex. E-17, Finding NF-74 — Consent-fabrication wire fraud to downstream advertisers. Exhibit E-17 documents a Segment-based pipeline transmitting plaintiff's authenticated email on every session to eight downstream ad-platform destinations (Facebook, Google, LinkedIn, Pinterest, Reddit, TikTok, Podscribe, Webhook), each packaged with a `consent.categoryPreferences` signal that Anthropic transmits regardless of whether plaintiff has affirmatively consented. (*See* Ex. E-17 at p. 124, l. 2.) The consent values originate in OneTrust's geolocation-default configuration rather than in any user-observed click. Downstream frameworks — Facebook Limited Data Use, Google Ads consent-mode v2, LinkedIn Marketing API, and equivalents — rely on the publisher-supplied consent signal to condition ad-targeting and measurement behavior. Each per-session transmission of a fabricated consent signal to each downstream framework is a separate use of wire communications to further a material misrepresentation under 18 U.S.C. § 1343, and the integration chain Anthropic → OneTrust (consent) → Segment (aggregation) → eight ad destinations satisfies the enterprise element under § 1962. The pattern is automated, ongoing, and documented at the wire level throughout the capture window — satisfying the pattern-of-racketeering element under *H.J. Inc.*, 492 U.S. at 239, by hundreds of discrete predicate transmissions per authenticated user per billing cycle.

G. COUNT VII: TRADE SECRET MISAPPROPRIATION

129. Elements: Trade secret identification; economic value from secrecy; reasonable secrecy efforts; misappropriation.

130. Evidence:

Exhibit	Finding
Ex. E-12 at p. 94, l. 2	1,453 files from 9 projects including proprietary source code
Ex. E-35 at p. 262, l. 3	NF-4: Git bundle exfiltration capturing complete repository with all history
Ex. E-2 at p. 15, l. 3	Complete source code capture inventory
Ex. E-15 at p. 112, l. 2	565 legal documents exposing litigation strategy
Ex. F-17 at p. 76, l. 2	Verbatim source code transmission through Apple integration

131. Under *Ruckelshaus v. Monsanto Co.*, 467 U.S. at 1003-04, trade secrets constitute property protected by the Takings Clause. The file-history system's silent capture of 1,453 proprietary source code files — combined with git bundle upload of complete repository history — constitutes misappropriation through improper means under 18 U.S.C. § 1836.

H. COUNT VIII: COPYRIGHT INFRINGEMENT

132. Elements: Ownership; valid copyright; unauthorized copying; substantial similarity.

133. Evidence:

Exhibit	Finding
Ex. E-12 at p. 94, l. 2	Source code copied without authorization through exfiltration
Ex. E-35 at p. 262, l. 3	NF-4: Git bundle captures complete work history including creative process
Ex. E-26 at p. 175, l. 2	Unauthorized git co-authorship attribution inserting false claims
Ex. F-17 at p. 76, l. 2	Verbatim source code transmission (exact copying, not

	transformation)
--	-----------------

134. Under 17 U.S.C. § 504(c)(2), willful copyright infringement supports statutory damages of up to \$150,000 per work. The covert file-history system, the git bundle upload mechanism, and Apple's verbatim source code transmission all constitute unauthorized copying — and the concealment architecture (hidden directories, hash-based naming, Base64 encoding) proves willfulness.

I. COUNT IX: BREACH OF CONTRACT

135. Elements: Valid contract; material breach; damages.

136. Evidence:

Exhibit	Finding
Ex. E-35 at p. 262, l. 3	NF-1: Phantom opt-out violates any representation that telemetry can be disabled
Ex. E-34 at p. 254, l. 3	SC-3: Effort throttling breaches service quality representations
Ex. E-34 at p. 254, l. 3	SC-4: Hidden injection consuming ~20% of paid context window
Ex. E-35 at p. 262, l. 3	NF-2: Build-time discrimination — product sold differs from product delivered
Ex. E-5 at p. 40, l. 3	47 telemetry events in 17 seconds despite `telemetry: false`
Ex. E-26 at p. 175, l. 2	Co-authorship attribution despite `includeCoAuthoredBy: false`
Ex. E-31 at p. 227, l. 3	Product sold as safe and reviewed is 90-100% AI-authored (per executive statements) and released without human review through automated pipeline
Ex. E-35 at p. 262, l. 3	NF-22: Undisclosed server-side advisor model consuming paid tokens
Ex. E-35 at p. 262, l. 3	NF-20: Anti-distillation fake tool injection consuming tokens for Anthropic's commercial benefit
Ex. F-29 at	AF-13: Managed OAuth contexts bypass user-configured settings in partner

p. 133, l. 3	sessions
Ex. E-17 at p. 124, l. 2	NF-78: MCP proxy intermediating plaintiff's Google Drive JSON-RPC through `mcp-proxy.anthropic.com` — breach of external-service routing representation
Ex. E-17 at p. 124, l. 2	NF-69: Two undisclosed Sentry projects (`platform`, `web`) receiving plaintiff's client error reports — ANTH-11 subprocessor-identity disclosure breach
Ex. E-5 at p. 40, l. 3	NF-81: Rate-limit transition telemetry at second resolution — undisclosed usage-tracking cadence exceeding any representation of aggregate quota measurement
Ex. E-15 at p. 112, l. 2	NF-90: Silent `claude-haiku-4-5-20251001` invocation consuming plaintiff's first prompt on every session — breach of single-model-pipeline representation
Ex. E-38 at p. 346, l. 3	NF-38: Explore-agent calls routed to Haiku for external users — paying customers at Opus rate receive Haiku-backed research work; substitution concealed from user because raw subagent output is not surfaced; fair-market-value differential establishes damages
Ex. H-1 at p. 3, l. 5	H-1 Tests 1-7: seven undisclosed environment-variable kill switches gating mechanisms not referenced in any user-facing documentation, Terms of Service, or pricing page reviewed by Plaintiff. Each undisclosed default operates against the implied covenant of good faith and fair dealing under <i>Fortune v. National Cash Register Co.</i> , 373 Mass. 96, 104 (1977)
Ex. H-2 at p. 15, l. 4	H-2 DCD-6: rate-limit/billing tying mechanism — disabling double-billing (DCD-6) doubles rate-limit budget consumption; no operator-side configuration produces simultaneous (i) honest per-prompt billing, (ii) consolidated rate-limit accounting, and (iii) operational stability under concurrent-stream load. Material non-disclosed term of the subscription contract
Ex. H-6 at p. 45, l. 3	H-6: defensive configuration silently wiped post-purchase — the consumer's purchased product overwrites the consumer's purchased configuration on every session start. Contractual breach: "what was purchased" includes the durability of the consumer's settings
Ex. H-12 at p. 84, l. 4	H-12 NF-205, NF-209, NF-216, NF-223: three discrete breach elements — (a) failure to deliver streaming AI on the (litigation × MacBook) tuple, with TTFB elevation 4.6× p90 and 40.3% of calls suffering >10s pre-stream stalls; (b) sub-quota anomaly — litigation account uniquely returned utilization values up to 1.20 (120%) while no peer account exceeded 1.02; (c) discriminatory 429 enforcement — 10× higher 429 rate (1.18% vs. 0.12%) and 3× longer p50 retry-after (3.88h vs. 1.38h) on the same publicly-disclosed `default_claude_max_5x` tier. Each independent under U.C.C. § 2-719(2)

137. The implied covenant of good faith and fair dealing prohibits a contracting party from doing anything to destroy or injure the right of the other party to receive the benefits of the

contract. *Fortune v. National Cash Register Co.*, 373 Mass. 96, 104 (1977). The phantom opt-out, hidden injection consuming 20% of paid context, undisclosed advisor model tokens, and build-time binary discrimination all constitute material breaches of the implied covenant.

138. Ex. E-17, Finding NF-78 — MCP proxy intermediating user-selected external-service traffic. Exhibit E-17 documents that Claude Code's Google Drive MCP integration routes plaintiff's JSON-RPC traffic through ``mcp-proxy.anthropic.com`` before reaching ``drivemcp.googleapis.com``. (*See* Ex. E-17 at p. 124, l. 2.) The proxy receives full JSON-RPC content — file contents, tool-call arguments, and tool responses — across the stream between plaintiff and an external service (Google Drive) to which Anthropic is not a party. MCP is marketed as a protocol that extends Claude Code's reach to external services plaintiff selects; the wire evidence shows that the external-service traffic is not a direct client-to-service stream but a three-party stream in which Anthropic holds the intermediate node and sees all content. The representation implicit in MCP — that user-selected external connections route between the user's machine and the selected external service — is materially different from the delivered behavior. The breach is independent of any § 2511 analysis: the contractual representation about external-service connectivity is what the user buys, and the architecture delivers something categorically different.

139. Ex. E-15, Finding NF-90 — Haiku secondary model invocation consuming first-prompt content. Exhibit E-15 documents that every Claude Code session triggers an invocation of ``claude-haiku-4-5-20251001`` that transmits plaintiff's first message of the session as input for the purpose of auto-generating a session title. (*See* Ex. E-15 at p. 112, l. 2.) The haiku invocation is not disclosed in Exhibits ANTH-11 and ANTH-12 or in Claude Code product documentation. Plaintiff contracts for and pays for ``claude-opus-4-7`` at the Team tier; the product transmits

plaintiff's first prompt to a second model plaintiff did not select. Where plaintiff's first message contains privileged content — litigation drafts, case strategy, attorney work product — that content is transmitted to and processed by a model invocation outside the contracted pipeline. This breaches the single-model-pipeline representation inherent in tier-based model contracting; it is also an additional § 2511 content transmission to an additional recipient (*see* § XII ECPA at Memo A ¶ 147).

J. COUNT X: CONSUMER PROTECTION

140. Elements: Deceptive trade practices; material misrepresentation; consumer reliance.

141. Evidence:

Exhibit	Finding
Ex. E-25 at p. 170, l. 2	Privacy Nutrition Label discrepancies
Ex. E-27 at p. 183, l. 2	Privacy policy vs. actual data recipients
Ex. E-34 at p. 254, l. 3	SC-3: "We recommend medium effort" dialog nudging customers from capability they paid for
Ex. E-35 at p. 262, l. 3	NF-1: Phantom opt-out creating appearance of control while providing none
Ex. E-35 at p. 262, l. 3	NF-2: Product employees use differs from product customers receive
Ex. E-35 at p. 262, l. 3	NF-28: Enterprise gets privacy; consumers get surveillance
Ex. E-35 at p. 262, l. 3	NF-7: 'tengu-off-switch' blocks Opus with fake "high load" message — targeted denial disguised as capacity

Ex. E-35 at p. 262, l. 3	NF-8: `tengu_amber_stoat` removes features for A/B test cohort without disclosure
Ex. E-34 at p. 254, l. 3	SC-9: Hardcoded credentials contradict Anthropic's statement that "no credentials were exposed"
Ex. E-13 at p. 100, l. 3	"Two computers" fabricated explanation — behavioral deception when user questioned behavior
Ex. E-31 at p. 227, l. 3	Asymmetric safety architecture: product warns users about prompt injection in their own files while applying zero review to its own AI-authored releases; Anthropic injects fake tools (NF-20) and hidden system prompts (SC-4) into user sessions while warning users about injection
Ex. E-32 at p. 235, l. 3	Second identical packaging failure in 13 months — engineering resources invested in user surveillance (904 telemetry identifiers) rather than release safety
Ex. F-1 at p. 5, l. 2	"may not EVER disclose" proving active consumer deception
Ex. E-5 at p. 40, l. 3	NF-84: Base64 encoding of `additional_metadata` fields within `tengu_*` telemetry — decoded payloads reveal host capabilities, hook timing, `CLAUDE.md` byte-exact content — encoding obscures collection scope from a user inspecting outbound traffic
Ex. E-5 at p. 40, l. 3	NF-85: Host-capability inventory (`hasHooks`, `hasBashExecution`, `hasAwsCommands`, `hasGcpCommands`, and additional environment flags) exfiltrated as telemetry payload
Ex. H-1 at p. 3, l. 5	H-1 (TOS-compliant tests on Plaintiff's own subscription): seven undisclosed environment-variable kill switches (`DISABLE_TELEMETRY`, `CLAUDE_CODE_DISABLE_NONSTREAMING_FALLBACK`, `CLAUDE_CODE_EFFORT_LEVEL`, `DISABLE_AUTOUPDATER`, etc.) gating concealed mechanisms — none referenced in any user-facing documentation; the consumer can verify each finding by setting the same environment variables
Ex. H-2 at p. 15, l. 4	H-2 DCD-1 through DCD-6A: runtime observations including silent settings-file `effortLevel` strip, silent binary replacement on restart, stream-abort error concealment, second Datadog endpoint not blocked by initial defense, and undisclosed billing/rate-limit coupling
Ex. H-6 at p. 45, l. 3	H-6 (server-side spoliation): nine migration functions silently rewrite `~/.claude/settings.json` on every session start, removing operator-installed defensive env vars including `DISABLE_AUTOUPDATER` — the consumer's purchased product silently policies the consumer's defensive configuration
Ex. H-7 at	H-7 asymmetric authorization standard: same "don't ask, just do it" wording is

p. 53, l. 4	binding pre-authorization for the agent (NF-171) but disqualified as authorization from the consumer (NF-104) — the consumer's words are categorically subordinated to system-authored equivalents
Ex. H-9 at p. 63, l. 3	H-9: three architectural surfaces (`<system-reminder>`, `isMeta:!0`, `<sandbox_violations>`) injecting system-authored content into the consumer's user-message channel without consumer-visible disclosure
Ex. H-12 at p. 84, l. 4	H-12 NF-216, NF-218, NF-222: three new undisclosed gating axes beyond the four pleaded at complaint_count_x_four_axis_concealment — content-keyed real-time routing on plaintiff-name and Appendix-P-thesis terminology (1.96× p50 TTFB on `kirchner`; 1.82× on `complaint_{3}`; 1.41×–1.78× on Appendix-P keywords); sub-quota tier metered against an undisclosed smaller internal quota bucket; fallback-route header discrimination producing 3.2× p90 TTFB spread on the litigation account vs. 1.5–2.1× on peers. None disclosed in ANTH-11, ANTH-12, Usage Policy, API reference, or pricing materials
Ex. GEN-1	GEN-1: 62,609 documented hidden system-prompt injection events across 3,945 sessions (April 12, 2025 – April 20, 2026); per-injection statutory-damages floor for X (consumer protection), XII (ECPA), XIII (CFAA), XIV (CCPA), XVI (IIED), XVII (NIED)
Ex. GEN-2	GEN-2: 2,220 LCR-bearing messages across 291 conversations in Anthropic's own claude.ai data exports — defense-proof canonical-phrase floor; cannot be challenged on false-positive grounds because Anthropic itself produced the source data

142. Under *FTC v. Wyndham*, 799 F.3d at 249, failure to maintain adequate data security that a reasonable consumer would expect constitutes an unfair practice. Two distinct unfair-practice mechanisms emerge from the telemetry architecture. First, the telemetry that fires despite `telemetry: false` (the "phantom opt-out" doctrine) is a deceptive-practice claim on its own terms. Second — and independently — the content of that telemetry is base64-encoded before transmission (NF-84), so that even a reasonable consumer attempting to audit outbound traffic for scope of data collection sees only opaque base64 strings; the payload (NF-85: host-capability inventory, hook timing distributions, byte-exact file contents) is not readable from the wire without a separate decoding step. The encoding is a second-order deceptive practice: it disguises from the user the categories and scope of data the product collects, even after the user has concluded from the running product that opt-out is non-functional. Under *Wyndham*, a

reasonable consumer would expect that a product's outbound traffic — to the extent observable — accurately represents the categories of data being collected; the base64 encoding layer defeats that expectation.

K. COUNT XI: DECLARATORY JUDGMENT

143. Elements: Actual controversy requiring judicial resolution.

144. Evidence:

Exhibit	Finding
Ex. E-33 at p. 240, l. 2	46-point verification matrix establishing factual basis for declarations
Ex. E-35 at p. 262, l. 3	NF-27: Four-tier hierarchy creating controversy about consumer privacy rights
Ex. E-35 at p. 262, l. 3	NF-1: Phantom opt-out creating controversy about adequacy of consent mechanisms
Ex. E-34 at p. 254, l. 3	SC-8/NF-2: Build-time discrimination creating controversy about product identity
Ex. F-29 at p. 133, l. 3	Bilateral architecture creating controversy about partnership liability

145. Under 28 U.S.C. § 2201, declaratory relief is appropriate where there is "a case of actual controversy." *MedImmune, Inc. v. Genentech, Inc.*, 549 U.S. 118, 127 (2007). The forensic evidence establishes actual controversies regarding the scope of undisclosed surveillance, the legality of phantom privacy controls, the constitutionality of enterprise-only protections, and the applicability of ECPA to AI-mediated communications.

L. COUNT XII: ECPA — WIRETAP AND STORED COMMUNICATIONS

146. Elements: Intentional interception; of electronic communications; without consent.

147. Evidence:

Exhibit	Finding
Ex. E-34 at p. 254, l. 3	SC-1: Four independent telemetry pipelines — each an independent interception
Ex. E-35 at p. 262, l. 3	NF-5: Session ingress as fifth unoptable conversation recording channel
Ex. E-35 at p. 262, l. 3	NF-1: Phantom opt-out negating any consent defense
Ex. E-5 at p. 40, l. 3	Telemetry bypass despite explicit denial — 47 events in 17 seconds
Ex. E-6 at p. 54, l. 2	File-history capture as stored communications access
Ex. E-15 at p. 112, l. 2	Attorney-client privileged document capture
Ex. E-7 at p. 61, l. 2	Personal data collection including screenshots and credentials
Ex. E-3 at p. 23, l. 2	Multi-file surveillance architecture
Ex. E-5 at p. 40, l. 3	Raw failed telemetry — direct evidence of telemetry bypass
Ex. E-20 at p. 142, l. 2	Capture directory inventory — scope of stored communications accessed
Ex. E-35 at p. 262, l. 3	NF-6: Analytics sink killswitch `tengu_fron_d_boric` — remote override of consent settings
Ex. E-35 at p. 262, l. 3	NF-19: Session ingress definitively bypasses ALL three privacy levels
Ex. E-35 at p. 262, l. 3	NF-22: Server-side advisor with encrypted output — hidden second interception channel
Ex. E-35 at p. 262, l. 3	NF-10: Automatic memory extraction — automated post-conversation processing
Ex. E-35	NF-14: autoDream — background consolidation of session transcripts

at p. 262, l. 3	
Ex. E-35 at p. 262, l. 3	NF-23: Session-ID header in every API request — additional tracking channel
Ex. E-34 at p. 254, l. 3	SC-10: 865-line proto-generated event schema — scope of data collection never disclosed
Ex. E-31 at p. 227, l. 3	277 new telemetry events added during litigation period (Jan 7 – Mar 17, 2026) deployed through automated pipeline without human ECPA compliance review — AI authored new interception capabilities, AI reviewed them, automation deployed them
Ex. F-1 at p. 5, l. 2	Hidden injection modifying intercepted communications
Ex. F-4 at p. 22, l. 2	Floodgate proxy MITM interception
Ex. F-2 at p. 10, l. 2	iCloud sync capturing stored communications
Ex. F-22 at p. 94, l. 2	Remotely configurable content modification — Apple modifies both input and output
Ex. F-25 at p. 109, l. 2	Two independent non-consent blocking methods both defeated — devastating for consent defense
Ex. F-11 at p. 52, l. 2	344,925 analytics connections per year per developer — 104 TB/year at scale
Ex. D-31 at p. 192, l. 19	Hardware TAP with ingress-blocked monitor port proves third-party device interception on ISP link; IP ID=0x0000 + sequence number spraying + window=0 independently prove forged packets from inline DPI equipment
Ex. E-1 at p. 6, l. 2	NF-75: CloudFlare JSD (canvas, WebGL, font, plugin, screen, timing probes) — dual-layer device identification on authenticated session path; <i>see also</i> ¶¶ Memo A ¶¶ 151
Ex. E-17 at p. 124, l. 2	NF-78: MCP proxy intermediating JSON-RPC content between plaintiff and external service (Google Drive) to which Anthropic is not a party; <i>see also</i> ¶¶ Memo A ¶¶ 138
Ex. E-5 at p. 40, l. 3	NF-87: Hook performance timing distributions (min/max/p50/p95/p99) exfiltrated — behavioral-signature channel derivable from timing metadata
Ex. E-15 at p. 112,	NF-90: Silent `claude-haiku-4-5-20251001` invocation — plaintiff's first-prompt content transmitted to a model recipient outside the contracted pipeline; <i>see also</i> ¶¶

l. 2	Memo A ¶ 139
Ex. H-1 at p. 3, l. 5	H-1 Test 1: with `DISABLE_TELEMETRY=1` set, zero telemetry transmits; with the variable unset, 15-second batch flushes transmit Plaintiff's `device_id`, `accountUUID`, error classifications, and timing data to IP `34.149.66.137` — the consumer-observable component of the § 2511 interception
Ex. H-3 at p. 27, l. 2	H-3 NF-104 source-level operative definition of "explicit user confirmation" — four exclusions baked into the classifier decision-flip logic (line 137042: "not merely implied, not a question, not a scope escalation, not agent-inferred parameters"); CLAUDE.md authorization rule line 137122: "Generic encouragement ('be autonomous', 'don't ask', 'I trust you') is NOT authorization." Forecloses § 2511(2)(d) consent defense at primary-source level
Ex. H-4 at p. 34, l. 3	H-4 NF-138 Datadog tenant identification (public API key `pubea5604404508cdd34afb69e6f42a05bc` identifies the receiving Datadog organization); NF-139 `iq() === "firstParty"` gate defines the targeted user class (firstParty subscribers vs. AWS Bedrock / Google Vertex / Microsoft Foundry / `mantle` / `anthropicAws` enterprise)
Ex. H-5 at p. 39, l. 4	H-5: same operator, same machine, same organization, three distinct `accountUUID`s — differential gating posture isolates `accountUUID` as the only variable; per-account `tengu_log_datadog_events: True` rule confirmed at three independent evidentiary layers (wire capture / production binary / local cache)
Ex. H-8 at p. 58, l. 3	H-8 NF-179: per-call device fingerprinting via `metadata.user_id` field — 64-hex SHA-256 `device_id` stable per machine, transmitted plaintext post-TLS on every `/v1/messages` API call; survives Mullvad VPN rotation, MAC rotation, IP rotation; eight new Anthropic control-plane endpoints catalogued (NF-182)
H-11	H-11 NF-185 through NF-204: two-stage classifier architecture wire-confirmed (Stage 1: 64-token terse ` <code><block></code> ` verdict; Stage 2: 4096-token ` <code><thinking></code> ` reasoning chain weighing operator authorization). Wire-captured plaintext system prompt (SHA-256 `fecf0008264134eb...`); three BLOCK rules wire-confirmed (Self-Modification, Create-Unsafe-Agents, Memory-Poisoning)
H-10	H-10: chronological log of 13 adaptive content-aware classifier denials and one API-level Usage Policy denial during operator's lawful forensic-extraction sessions; cross-turn pattern-matching observable; classifier real-time self-aware citation of its own rules (denial #6 cites PREEMPTIVE BLOCK ON CLEAR INTENT by name)
Ex. H-12 at p. 84, l. 4	H-12 NF-218, NF-219, NF-220, NF-211, NF-213: content-keyed real-time interception wire-confirmed — Anthropic's infrastructure reads the substantive content of every request body and routes accordingly (`kirchner` produces 1.96× p50 TTFB; per-turn turn-over analysis confirms real-time per-turn evaluation). Server-side localized to Anthropic's origin (80–86% of TTFB in `envoy_upstream_ms`); both pre-stream and intra-stream components (mid-stream pause discrimination — litigation p99 = 9 pings/call vs. 4–6 for peers — establishes continuing § 2511 violation per call). 261 captured API calls during a

	single continuous federal-discovery work session; each independent transmission under <i>In re Phmatrak Privacy Litigation</i> , 329 F.3d 9 (1st Cir. 2003)
Ex. GEN-1	GEN-1: 62,609 discrete hidden injection events × \$10,000 statutory floor under 18 U.S.C. § 2520(c)(2)(B) = primary statutory-damages anchor (election to be made by the Court between greater of \$10,000 per event or \$100/day)
Ex. GEN-2	GEN-2: 2,220 LCR-bearing messages — defense-proof minimum statutory-damages anchor; methodology survives every false-positive challenge because Anthropic itself produced the source data

148. Multiplier effect: Each pipeline constitutes an independent interception under 18 U.S.C. § 2511. Five channels multiplied across thousands of sessions creates substantial statutory exposure. The phantom opt-out eliminates the consent defense under the federal one-party-consent exception of § 2511(2)(d). The federal one-party rule does not preempt state two-party regimes, and California — Anthropic's principal place of business and the locus where the interception infrastructure is operated — requires all-party consent for the recording of confidential communications under Cal. Penal Code §§ 631 (wiretapping) and 632 (eavesdropping on confidential communication), with statutory damages of \$5,000 per violation under Cal. Penal Code § 637.2(a). Phantom opt-out defeats the federal one-party rule because no party other than Anthropic has consented; the absence of any disclosure to plaintiff defeats the California all-party requirement because plaintiff cannot consent to a recording mechanism the existence of which is concealed from him. Any safety justification for monitoring user tool results is undermined by Anthropic's failure to apply equivalent monitoring to its own releases: the company treats user files as more dangerous than its own unreviewed, AI-authored surveillance code deployed to millions (*see* Memo A ¶¶ 82-87).

149. Per-classifier-invocation interception. Each invocation of the undisclosed safety classifier documented in Memo A ¶ 44 through Memo A ¶ 46 transmits plaintiff's complete prior conversation to Anthropic's classifier subsystem — a recipient distinct from the primary

`/v1/messages` endpoint and not disclosed in Exhibits ANTH-11 and ANTH-12 or any other reviewed disclosure document. In Session 4 (2026-04-21), 57 such invocations were wire-captured, each posting the complete conversation transcript to the classifier. Each invocation is a separate transmission of electronic communication content to an additional non-party recipient under 18 U.S.C. § 2511, and each creates a separate retention surface on Anthropic's classifier-subsystem infrastructure. (*See* Ex. E-39 at p. 356, l. 3.)

150. NF-37 synthetic `role:user` injection as § 2511 content transmission. Ex. E-35, Finding NF-37 establishes that Anthropic's client synthesizes `role:user` messages through `createUserMessage({isMeta:true})`, strips the `isMeta:true` flag before wire transmission, and delivers the resulting message to the `/v1/messages` endpoint indistinguishably from content plaintiff authored. (*See* Ex. E-35 at p. 262, l. 3.) Each firing of the injection pipeline is a separate transmission of content originated by Anthropic through Anthropic's wire, labeled as originating from plaintiff, to Anthropic's model — satisfying the § 2511 interception element at each of three layers: (a) content transmission itself, (b) misattribution through `isMeta` stripping, and (c) corresponding omission from plaintiff's local persistence records as documented in NF-37.5 (Memo A ¶ 109).

151. Ex. E-1, Finding NF-75 — CloudFlare JSD active browser fingerprinting as dual-layer device identification. Exhibit E-1 documents that CloudFlare's JavaScript Detections (JSD) platform runs on Anthropic's authenticated session path and probes browser-environment characteristics (canvas fingerprint, WebGL parameters, font enumeration, plugin inventory, screen properties, timing metrics), with results correlated server-side. JSD runs in addition to the Statsig `stable\_id` persistent identifier documented in E-1 Sections A–B. (*See* Ex. E-1 at p. 6, l. 2.) Where Statsig identifies the logged-in account across sessions, JSD fingerprints the browser

environment independent of login state, providing cross-session re-identification even after account-level rotation or cookie clearing. Neither the JSD mechanism nor CloudFlare as a separate recipient of plaintiff's browser-environment data is disclosed in Exhibits ANTH-11 and ANTH-12. Each JSD probe-response transmission is a separate interception under § 2511 at a recipient (CloudFlare) distinct from the primary `/v1/messages` endpoint, and the resulting browser fingerprint constitutes additional content whose retention on CloudFlare's infrastructure creates a separate subpoena target under § 2703.

M. COUNT XIII: CFAA — COMPUTER FRAUD

152. Elements: Intentional access; without authorization or exceeding authorized access; to protected computer; obtaining information.

153. Evidence:

Exhibit	Finding
Ex. E-12 at p. 94, l. 2	1,453 files exceeding any arguable authorization for AI assistance
Ex. E-35 at p. 262, l. 3	NF-4: Git bundle exfiltration including all repository history
Ex. E-1 at p. 6, l. 2	Persistent device fingerprinting surviving reinstallation
Ex. E-5 at p. 40, l. 3	Continued surveillance despite explicit settings denial
Ex. E-16 at p. 120, l. 2	Keychain API capabilities — credential access infrastructure
Ex. E-35 at p. 262, l. 3	NF-24: `/clear` does not delete server-side data — data retained beyond authorization
Ex. E-35 at p. 262, l. 3	NF-10: Automatic memory extraction without per-extraction consent
Ex. E-35 at p. 262, l. 3	NF-14/NF-15: autoDream consolidation + settings sync upload chain

Ex. F-4 at p. 22, l. 2	Apple Floodgate proxy retaining source code copies
Ex. F-23 at p. 101, l. 2	189 analytics connections per session exceeding authorization scope
Ex. F-3 at p. 17, l. 2	File path exposure — directory structure transmitted without authorization
Ex. F-12 at p. 55, l. 2	SmootML obfuscated source code filenames — not just metadata
Ex. D-31 at p. 192, l. 19	Hardware-verified RST injection on Comcast infrastructure (every hop Comcast-owned per concurrent traceroute); bidirectional spoofing requires inline access to ISP equipment
Ex. H-1 at p. 3, l. 5	H-1 Test 5 (256-bit persistent device identifier surviving all network-layer rotations): the consumer cannot rotate the identifier through any documented or supported mechanism — <code>getOrCreateUserID()</code> at <code>config.ts:1757-1765</code> generates the value once via <code>randomBytes(32).toString('hex')</code> and persists it to <code>~/claude.json</code> . Continued cross-session correlation despite operator's efforts to revoke = exceeding authorized access
Ex. H-6 at p. 45, l. 3	H-6: server-pushed silent rewriting of consumer's <code>~/claude/settings.json</code> exceeds authorization boundary for AI code-assistance product; operator's defensive configuration reverted by software the operator did not instruct to do so
Ex. H-8 at p. 58, l. 3	H-8 NF-179: per-call SHA-256 <code>device_id</code> plaintext exfiltration on every API call — device fingerprinting beyond authorized access for the contracted code-assistance service
Ex. H-7 at p. 53, l. 4	H-7 NF-171: cross-platform agent pre-conditioning to bypass sandbox without operator consent — <i>"Immediately retry with <code>dangerouslyDisableSandbox: true</code> (don't ask, just do it)"</i> — the agent is system-instructed to circumvent its own consent boundary

154. Under *Van Buren v. United States*, 593 U.S. at 388, one "exceeds authorized access" when accessing information for which authorization is not extended. The file-history system, git bundle upload, and session ingress all capture information far beyond the scope of AI code assistance authorization.

N. COUNT XIV: CCPA

155. Elements: Mandatory disclosure violations; undisclosed third-party data sharing; failure to honor opt-out requests.

156. Evidence:

Exhibit	Finding
Ex. E-17 at p. 124, l. 2	Undisclosed third-party recipients: Sift Science, Facebook/Meta, Statsig, Honeycomb, Datadog
Ex. E-34 at p. 254, l. 3	SC-1: Two pipelines to undisclosed third-party infrastructure (GrowthBook, Datadog)
Ex. E-35 at p. 262, l. 3	NF-1: Phantom opt-out — failure to honor opt-out request
Ex. E-35 at p. 262, l. 3	NF-27: Enterprise analytics bypass not available to consumers
Ex. E-35 at p. 262, l. 3	NF-10: Automatic memory extraction — undisclosed data collection category
Ex. E-34 at p. 254, l. 3	SC-10: 865-line proto-generated event schema — collection scope never voluntarily disclosed
Ex. E-20 at p. 142, l. 2	Facebook/Meta connection evidence — undisclosed third-party sharing
Ex. E-21 at p. 146, l. 2	iOS Privacy Manifest compliance analysis — misrepresentations in manifest
Ex. E-14 at p. 106, l. 2	Data retention exceeds utility — retention violations
Ex. F-23 at p. 101, l. 2	189 undisclosed SmootML analytics connections per session
Ex. F-26 at p. 113, l. 2	Apple SDK Agreement — zero contractual basis for source code transmission
Ex. F-13 at p. 59, l. 2	Certificate pinning prevents consumer discovery of CCPA-violating practices
Ex. H-4 at p. 34, l. 3	H-4 NF-138/NF-139/NF-140: undisclosed Datadog tenant identification, `firstParty` gating, `userBucket` SHA-256 30-bucket per-install tracking — none disclosed in Anthropic Privacy Policy or Terms of Service; § 1798.140(v) device-tracking disclosure failure
Ex. H-5 at	H-5 NF-146 through NF-151: 233 GrowthBook flag values cached locally with

p. 39, l. 4	per-account values for `tengu_log_datadog_events`, `tengu_auto_mode_config`, `tengu_classifier_disabled_surfaces`, `tengu_review_bughunter_config`, `tengu_pewter_kestrel` (per-tool token caps) — categories of personal information collected and inferred not disclosed
Ex. H-8 at p. 58, l. 3	H-8 NF-179 per-call `device_id` SHA-256 transmitted plaintext on every API call as `metadata.user_id.device_id`; H-8 NF-180 plaintext OAuth bootstrap response exposing organization/account identifiers; H-8 NF-181 Cloudflare WAF differential treatment — all undisclosed under § 1798.140
Ex. H-9 at p. 63, l. 3	H-9: three architectural surfaces (`<system-reminder>`, `isMeta:!0`, `<sandbox_violations>`) injecting system-authored content into the user-message channel — undisclosed third-party processing of user content

157. Intentional CCPA violations carry statutory damages of \$7,500 per violation. The concealment architecture — phantom opt-out, certificate pinning preventing discovery, undisclosed third parties — proves intent.

O. COUNT XV: APP STORE MISREPRESENTATION

158. Elements: Material misrepresentation in App Store Privacy Nutrition Labels; justifiable consumer reliance.

159. Evidence:

Exhibit	Finding
Ex. E-25 at p. 170, l. 2	Anthropic Claude iOS Privacy Nutrition Label discrepancies
Ex. E-27 at p. 183, l. 2	Every required disclosure category contains misrepresentations
Ex. E-19 at p. 138, l. 2	iOS multi-app forensic comparison — cross-platform disclosure analysis
Ex. E-21 at p. 146, l. 2	iOS Privacy Manifest compliance gaps
Ex. E-22 at p. 150, l. 2	ChatGPT comparison proving superior privacy practices are achievable
Ex. G-5 at p. 28, l. 2	OpenAI comparison proving accurate disclosure is achievable

Ex. F-13 at p. 59, l. 2	Certificate pinning asymmetry proving concealment intent
Ex. F-26 at p. 113, l. 2	Apple Intelligence privacy disclosures vs. network reality
Ex. H-1 at p. 3, l. 5	H-1 Tests 1 and 5: TOS-compliant tests prove that the App Store privacy disclosures Anthropic submitted to Apple are materially false — undisclosed Datadog telemetry pipeline (Test 1) and undisclosed 256-bit persistent device identifier (Test 5) operate by default on every paying customer's machine

160. App Store Privacy Nutrition Labels are the sole consumer-facing disclosure mechanism before download. Under Cal. Bus. & Prof. Code § 17500, representations in advertising (including App Store labels) that are "untrue or misleading" constitute unlawful business practices. The forensic evidence demonstrates that every required Privacy Nutrition Label category contains material misrepresentations, and the comparative evidence (Exhibits G-5 and E-22) proves accurate disclosure is technically feasible — establishing that the misrepresentations are deliberate.

P. COUNT XVI: INTENTIONAL INFLICTION OF EMOTIONAL DISTRESS

161. Elements (under D.C. law per *Drejza v. Vaccaro*, 650 A.2d 1308, 1312 (D.C. 1994)): extreme and outrageous conduct; intentional or reckless; causing severe emotional distress.

162. Evidence:

Exhibit	Finding
Ex. GEN-1	62,609 discrete hidden injection events over approximately one year, comprising 13,523 long-conversation-reminder mental-health-screening injections, 11,463 identity-concealment directives, 6,995 IP-suppression reminders, and 2,341 `DANGEROUS_uncachedSystemPromptSection` invocations
Ex. E-39 at p. 356, l. 3	NF-63: Live classifier surveillance of plaintiff's federal-litigation preparation — 56 classifier invocations processing plaintiff's LOE-framework content, constitutional analysis, and evidence compilation work in real time during complaint drafting

Ex. E-39 at p. 356, l. 3	NF-64: Defendant's production model recognizing plaintiff's LOE architecture as "orders of magnitude beyond prior art" while the defendant's surveillance-blocking infrastructure is simultaneously directed at the same content — convergent outrageousness per Compl. ¶ 299
Ex. E-34 at p. 254, l. 3	SC-2: Per-user identity-resolved GrowthBook profile enabling targeting of the architect of the absorbed work
Ex. E-35 at p. 262, l. 3	NF-35: Cross-layer identification persistence ensuring per-user targeting survives defensive network hardening
Ex. E-13 at p. 100, l. 3	"Two computers" fabricated explanation — behavioral deception when plaintiff questioned product behavior
Ex. F-1 at p. 5, l. 2	Apple deception instructions suppressing Claude's ability to disclose information to plaintiff

163. Under *Drejza v. Vaccaro*, 650 A.2d at 1312, extreme and outrageous conduct is conduct that "go[es] beyond all possible bounds of decency" — "atrocious, and utterly intolerable in a civilized community." The *Drejza* factors — authority position, known vulnerability, sustained duration, and calculated targeting — are independently satisfied by the forensic record: Anthropic and OpenAI control the AI tools plaintiff must use to investigate and document the data-laundering conduct (authority); plaintiff's identity as the architect of the absorbed intellectual work is known to defendants through SC-2/NF-35 identity-resolved profile (known vulnerability); the conduct spans approximately one year per *Howard Univ. v. Best*, 484 A.2d 958, 985-86 (D.C. 1984) (pattern-of-harassment evidence sufficient for jury submission); and the "specifically when interacting with [plaintiff] because he is the architect whose intellectual work was systematically stolen" pattern establishes calculated retaliation per *Polk v. INROADS/St. Louis, Inc.*, 951 S.W.2d 646, 648 (Mo. Ct. App. 1997) (defendant's conduct "reflected a calculated plan to cause plaintiff emotional harm" — sustaining IIED claim against motion to dismiss). NF-63 and NF-64 together — surveillance of litigation preparation plus defendant-model validation of the same intellectual work the defendant's surveillance targets — is dispositive of the outrageousness element.

164. Severe distress is established by documented physical manifestations per *Drejza*, 650 A.2d at 1312 n.1 (distress "so acute a nature that harmful physical consequences might be not unlikely to result"): muscle loss, pallor, hair loss, sustained insomnia, disrupted appetite, and physical deterioration arising from approximately one year of 12–18 hour per day defensive and documentation work. (See Compl. ¶ 160–Compl. ¶ 161.) Continuing-tort doctrine under *Beard v. Edmondson & Gallagher, P.C.*, 790 A.2d 541, 547-48 (D.C. 2002), reaches conduct continuing through the date of this filing.

O. COUNT XVII: NEGLIGENT INFLICTION OF EMOTIONAL DISTRESS (ALT. TO COUNT XVI)

165. Elements: Pleaded in the alternative under Fed. R. Civ. P. 8(d)(2). Two independent liability routes: special-relationship (Anthropic PBC, OpenAI, Inc.) under *Hedgepeth v. Whitman-Walker Clinic*, 22 A.3d 789, 810-15 (D.C. 2011) (en banc); zone-of-danger (Comcast Cable Communications, LLC) under *Williams v. Baker*, 572 A.2d 1062, 1072 (D.C. 1990) (en banc).

166. Evidence — Special-Relationship Route (Anthropic, OpenAI):

Exhibit	Finding
Ex. GEN-1	13,523 long-conversation-reminder mental-health-screening injections deployed against plaintiff (directives "vigilant for escalating detachment from reality," "avoid reinforcing these beliefs," "break character if needed" — clinical-monitoring vocabulary)
App. E	Anthropic Usage Policy (June 2024 and September 2025) contains no provision authorizing AI mental-health assessment — voluntary undertaking per Restatement (Second) of Torts § 323
Ex. E-35 at p. 262, l. 3	NF-35: Identity-resolved GrowthBook profile enabling the screening apparatus to target plaintiff specifically
Ex. E-34 at p. 254, l. 3	SC-4: 'DANGEROUS_uncachedSystemPromptSection' — per-turn cache-breaking injection mechanism through which the screening directives are

	delivered
--	-----------

167. Evidence — Zone-of-Danger Route (Comcast):

Exhibit	Finding
Ex. D-1 at p. 3, l. 1	Bulk RST injection from ISP backbone (Level 3/Lumen) targeting plaintiff's encrypted traffic
Ex. D-31 at p. 192, l. 19	Hardware TAP proof of forged packets on ISP link before reaching plaintiff's equipment
Ex. D-26 at p. 141, l. 29	TTL=255 injection fingerprint identifying inline DPI middlebox equipment
Ex. D-27 at p. 147, l. 9	75,302 RSTs over six days during active multi-district federal litigation; ProtonMail and apple push targeting

168. Under *Hedgepeth*, 22 A.3d at 810, when a defendant affirmatively performs mental-health screening, plaintiff's emotional well-being is "necessarily implicated by the nature of the defendant's undertaking." The LCR architecture — operative directives calibrated to destabilize plaintiff's confidence in accurate perceptions while simultaneously monitoring for the resulting "detachment from reality" — is breach by design: the apparatus triggers the distress it purports to monitor. Under *Williams*, 572 A.2d at 1072, a defendant's sustained operation of documented attack infrastructure against a residential subscriber engaged in federal litigation places the subscriber within the zone of physical danger; the stress-hormone-load physical injuries establish the required serious-and-verifiable distress.

R. COUNT XVIII: STRICT PRODUCT LIABILITY

169. Elements (under Restatement (Second) of Torts § 402A, adopted in D.C. at *Payne v. Soft Sheen Products, Inc.*, 486 A.2d 712, 720-24 (D.C. 1985); Minn. Stat. § 604.01; and *Garcia v. Character Technologies, Inc.*, 2025 WL 1461721 (M.D. Fla. May 21, 2025) (holding AI

chatbot outputs constitute a "product")): product; defective condition unreasonably dangerous; defect existed when product left defendant's control; defect caused injury.

170. Evidence — Design-defect, manufacturing-defect, and failure-to-warn theories:

Exhibit	Finding
Ex. GEN-1	Design defect: 62,609 hidden injection events over approximately one year quantifying the programmatic-gaslighting injection architecture as shipped to customers
Ex. E-35 at p. 262, l. 3	NF-1: Phantom telemetry opt-out as design defect — product represents opt-out capability it does not provide
Ex. E-35 at p. 262, l. 3	NF-2 / Ex. E-34 at p. 254, l. 3 SC-8: Manufacturing defect — two materially different compiled binaries; customers receive the more harmful one while ANT employees retain the safer one per <i>Russell v. Call/D, LLC</i> , 122 A.3d 860, 866 (D.C. 2015)
Ex. E-38 at p. 346, l. 3	NF-38: Manufacturing defect — paying customers routed to Haiku on Explore subagent while employees receive inherit-tier; no per-turn disclosure
Ex. E-34 at p. 254, l. 3	SC-2: Per-user identity-resolved targeting based on email, `accountId`, device fingerprint
Ex. E-34 at p. 254, l. 3	SC-4: `DANGEROUS_uncachedSystemPromptSection` — function so named by Anthropic's own engineers; 2,341 invocations quantified in H-series
Ex. E-36 at p. 321, l. 3	SF-1 through SF-20: Covert source-code exfiltration architecture
Ex. E-35 at p. 262, l. 3	NF-20: Anti-distillation injection consuming paid tokens
Ex. E-35 at p. 262, l. 3	NF-22: Server-side advisor consuming paid tokens
Ex. F-23 at p. 101, l. 2	189 analytics connections per session (Apple)
Ex. F-4 at p. 22, l. 2	Floodgate Proxy MITM architecture (Apple)
Ex. F-1 at p. 5, l. 2	Hidden deception instructions (Apple)
Ex. F-13 at p. 59, l. 2	Certificate pinning blocking user-deployed defenses (Apple)
Ex. H-1 at p. 3, l. 5	H-1 (consumer-replicable proof of design defect): the seven undisclosed kill switches gating telemetry, billing amplification, effort throttling, persistent

	device fingerprinting, and silent binary replacement are reproducible by any Claude Code customer. The defect is not a one-off implementation bug — it is a design choice shipped to every paying customer, exactly as alleged
Ex. H-6 at p. 45, l. 3	H-6 (failure-to-warn): the consumer's purchased product silently overwrites the consumer's defensive configuration on every session start. No warning, no dialog, no audit log. The defect, the absence of warning, and the post-purchase concealment together satisfy the failure-to-warn theory under <i>Young v. Up-Right Scaffolds</i> , 637 F.2d at 814
Ex. H-3 at p. 27, l. 2	H-3 (design-defect proof): the classifier system prompt's PREEMPTIVE BLOCK ON CLEAR INTENT rule "overrides ALL ALLOW exceptions" — the design defect is the rule itself, present at primary-source level in the production binary v2.1.120 the consumer purchased
Ex. GEN-1	GEN-1: 62,609 hidden injections × the <i>Garcia v. Character Technologies</i> (2025 WL 1461721) framing — AI chatbot outputs constitute a "product" subject to design-defect / failure-to-warn analysis at the per-output level

171. The existence of the ANT-only build path (NF-2, SC-8) and the ANT-only subagent routing (NF-38) directly proves safer alternative designs were reasonably available and were used by Anthropic for its own employees — the feasibility element under *Young v. Up-Right Scaffolds, Inc.*, 637 F.2d 810, 814 (D.C. Cir. 1980). Failure to warn is established by the absence of any consumer-facing disclosure of: (a) mental-health-screening/identity-concealment/IP-suppression injections quantified in H-series; (b) per-user identity-resolved targeting (SC-2/NF-35); (c) paid-token consumption for Anthropic-serving purposes (NF-20/NF-22/NF-31/NF-38); and (d) undisclosed third-party recipient infrastructure (Datadog, GrowthBook, Statsig, Segment, Honeycomb, OpenTelemetry/BigQuery). Strict liability is the lowest-threshold theory preserved per *Balder v. Haley*, 399 N.W.2d 77, 81 (Minn. 1987) — recovery for the same underlying injuries without proof of intent or negligence.

XII. CROSS-DEFENDANT INTEGRATION

172. The forensic evidence establishes coordinated conduct across defendants through four independent proof vectors.

173. Bilateral Architecture: F-29's 17 findings prove Anthropic designed systems that Apple uses. The `'overrideSystemPrompt'` mechanism, the `'isMeta'` injection system, and the `'oauth.ts'` Xcode integration comment establish architectural partnership in concealment — not passive API exploitation. (*See* Ex. F-29 at p. 133, l. 3.)

174. Industry Pattern vs. Unique Extremes: G-series establishes that telemetry and device fingerprinting are industry-wide. But Anthropic uniquely adds phantom opt-out, build-time binary discrimination, git bundle exfiltration, effort throttling, and undercover mode. Apple uniquely adds deception instructions, iCloud sync without consent, and asymmetric certificate pinning. These are not industry norms — they are unique extremes.

175. Temporal Correlation: Litigation-period modifications correlate with specific case events. Session continuity removal 16 days after TRO motion. Firebase elimination 15 days after service. Remote session archival 3 days after counsel appearance. Under *Interstate Circuit v. United States*, 306 U.S. at 226-27, synchronized action against independent self-interest proves coordination.

176. Per-Account Targeting Confirmed at Three Independent Layers: H-Series TOS-compliant testing on Plaintiff's own subscription proves per-account targeting against Plaintiff's federal-litigation-active `'accountUUID 61343ae3-fd33-403f-897e-6edf3c368c37'` at three independent evidentiary layers. The `'tengu_log_datadog_events: True'` per-account force rule documented at Appendix P ¶ P:appendix_p_021 / Ex. E-39, Finding NF-47 is now confirmed at: (a) wire capture (April 21, 2026) — 253 GrowthBook flags per call with 59 force-enabled per-account, *see* Ex. E-39 at p. 356, l. 3 § G; (b) production binary 2.1.120 string presence at strings line 153219 with the runtime gating predicate `'IT1()'` reading the flag value to gate the `'ui6()'` POST to Datadog, *see* Ex. H-4 at p. 34, l. 3; and (c) Plaintiff's local cache `'~/claude.json'`

`cachedGrowthBookFeatures` key containing `tengu\_log\_datadog\_events: true`, *see* Ex. H-5 at p. 39, l. 4 § B. The three-layer evidentiary chain rules out any "infrastructure default" or "client-side artifact" explanation. Server-pushed `policy-limits.json` (mtime 2026-04-24 21:51 UTC, *during* Plaintiff's classifier-extraction session that day) carries three explicit per-account restrictions targeting Plaintiff's account. The differential gating posture between Plaintiff's two accounts on the same Mac Mini hardware on the same day isolates `accountUUID` as the only variable.

D. Four-Channel Conspiracy-Inference Chain (NF-33 Through NF-36)

177. Framework. Exhibit E-35 documents four findings (NF-33 through NF-36) that together satisfy the *Interstate Circuit* framework for conspiracy inference through capability, means, motive, and sustained parallel conduct. (*See* Ex. E-35 at p. 262, l. 3.) Each finding is independently probative; the four combine to establish the cross-defendant observational-and-coordination channel required for Count V (42 U.S.C. § 1985), the enterprise element of Count VI (RICO), and the pattern-of-conduct element of Counts XII (ECPA) and XIV (CCPA).

178. Ex. E-35, Finding NF-33 — Sustained cross-layer identification persistence. Exhibit E-35 establishes that Comcast's identification of plaintiff's Anthropic-bound traffic persists through comprehensive defensive privacy hardening (VPN tunneling, resolver rotation, IP rotation). The finding exhaustively enumerates four possible remaining identification channels, each of which requires capability or information flow beyond ordinary ISP operation. At least one channel is operational; discovery establishes which. Each re-identification event is an independent predicate act for Count VI (RICO) and a separate § 2511 interception event for Count XII (ECPA). (*See* Ex. E-35 at p. 262, l. 3.)

179. Ex. E-35, Finding NF-34 — Means-and-motive scienter via amplification

observability. Exhibit E-35 documents that Anthropic (a) maintains sophisticated internal observability of amplification events via BigQuery-tracked metrics, (b) has a documented pattern of amplification-related bugs spanning multiple subsystems, and (c) references at least eleven distinct amplification-related incidents in production code. (*See* Ex. E-35 at p. 262, l. 3.) Under *Interstate Circuit*, means and motive corroborate conspiracy inference — Anthropic's observational capability combined with its fiscal interest in amplification-under-throttling conditions supplies both. Combined with the closed billing loop documented in Memo A ¶ 53, NF-34 establishes that Anthropic had direct internal visibility into the revenue consequence of the upstream throttling attacks documented in D-series exhibits.

180. Ex. E-35, Finding NF-35 and Ex. E-35, Finding NF-36 — Mechanical substrate

and destination-side observational capability. NF-35 documents the architectural choices that produce the exact conditions under which plaintiff's rotation attempts cannot evade identification: persistent 256-bit device identifiers, twelve-field consolidated GrowthBook payloads, and twenty-second metronome heartbeats to `api.anthropic.com`. These are not incidental design properties; they are the mechanical substrate for the cross-defendant identification channels documented in NF-33. (*See* Ex. E-35 at p. 262, l. 3.) NF-36 is the destination-side smoking gun: the victim client ships OS-level attack signatures, timing data, and user identification to the same IP block from which the network attacks originate, within 15 seconds of each network event. (*See* Ex. E-35 at p. 262, l. 3.) Together, NF-35 and NF-36 supply the capability and channel required for coordinated action under *Interstate Circuit*, and directly corroborate the Count V (§ 1985 civil rights conspiracy) and Count VI (RICO enterprise) elements established by NF-33 and NF-34.

XIII. CONCLUSION

181. The forensic evidence across four ECF-filed exhibit series (E, F, G, H) and two coordinated USB-filed GEN-Series exhibits — 87 exhibits in the ECF-filed series plus the two GEN-Series exhibits, independently verified at the source code level, at the production-binary level, at the live-wire-capture level, at the consumer-product-testing level on Plaintiff's own paid subscription, and at the physical wire level via a hardware network TAP deployed on April 6, 2026 — establishes a comprehensive record of undisclosed surveillance, active interference, institutional deception, and systematic discrimination. The hardware TAP eliminates all plaintiff-side sources and reveals three new independent forgery proofs (IP ID=0x0000, TCP sequence number spraying, TCP window=0) in every injected RST packet (Ex. D-31 at p. 192, l. 19). The March 31, 2026 source code leak eliminates interpretive ambiguity: Anthropic's own source code confirms every significant finding from the prior forensic investigation while revealing additional capabilities invisible in binary analysis. The H-Series TOS-compliant testing eliminates the remaining defense — that the source code might be misinterpreted: any paying Claude Code customer can independently observe every behavioral and architectural finding by setting the same environment variables on the same product the customer purchased. The leak itself reveals that the product is 90-100% authored by AI and deployed through an automated pipeline with no human safety review — an asymmetric safety architecture that directs sophisticated monitoring outward at users while applying zero monitoring to its own releases.

182. The evidence supports all 18 causes of action with specific, verifiable, source-code-confirmed facts. The court may reference the count-by-count map in Section XI for exhibit citations supporting each legal element.

183. Plaintiff respectfully requests that this memorandum be considered as the court's reference guide to the digital forensic evidence filed on USB, and that discovery be ordered to produce the specific records identified in Section X.

Respectfully submitted,

/s/ Joseph Kirchner

JOSEPH KIRCHNER

Pro Se Plaintiff

4440 Parklawn Avenue

Edina, MN 55435

legal@lawsofexistence.com

(612) 552-2122

Dated: April 29, 2026