

Claude for Elite Web Design — Research Report (2026-07-18)

For: Bob, Niven, and Paul

Recurring Research Report · definition: `Recurring-Claude-Web-Design-Research-Report.md` · web-researched + cited · prior run: `claude-web-design-2026-07-16.md`

Executive summary

Two days on from the 07-16 baseline, the fundamentals are unchanged — the field still diagnoses "distributional convergence" (AI slop) and still fixes it with compact taste-context up front plus a separate, seeing evaluator on the back end — but four developments this cycle are directly actionable for EdgePointe. First, **the model floor moved**: Opus 4.8 needs *less* design prompting to avoid slop, and its new `effort` parameter is load-bearing (it strictly scopes work at low/medium), while Fable 5 now leads WebDev Arena at 1653 Elo vs Opus 4.8's 1561 — so model + effort selection is now a design-quality lever, not just a cost lever. Second, **Impeccable** (Paul Bakaus) has emerged as the tool that most closely matches — and can plug straight into — our lint layer: 44 *deterministic, no-LLM* anti-slop rules plus ~17–23 invocable commands (audit/critique/polish/simplify/animate) that inherit your existing tokens instead of inventing a parallel system. Third, **Dynamic Workflows** (GA-ish JS orchestration, up to 16 concurrent / 1,000 total subagents with built-in adversarial cross-review) gives our mirror-score + parallel-builder pattern a native runtime. Fourth, Anthropic's official **frontend-design** skill has nearly doubled installs (277K → 565K) and now explicitly bans Space Grotesk, with an "Elite frontend-design-2026" variant carrying more aggressive typography/layout rules. Net: our genesis pipeline remains architecturally ahead; the highest-value new moves are (a) adopting Impeccable's deterministic detector as a fast pre-lint, (b) codifying model+effort selection per pipeline stage, and (c) porting the mirror-score loop onto Dynamic Workflows.

Findings

1. NEW — Model floor rose: Opus 4.8 needs less anti-slop prompting, and effort is now a design lever.

What it is: Anthropic's Opus 4.8 prompting guide states the model "requires less frontend design prompting than previous models to avoid generic AI-slop patterns" and generates distinctive frontends with minimal guidance. Critically, its `effort` parameter is respected strictly, especially at the low end: at low/medium the model "scopes its work to what was asked rather than going above and beyond," and Anthropic says effort matters more for this model than any prior Opus.

Why it matters: Under-setting effort will silently *cap* design ambition — a genesis draft run at low effort will look deliberately restrained, not elite. Conversely, heavy negative-prompt scaffolding that we wrote for older models may now be partially redundant overhead.

How EdgePointe applies it: Set **high effort explicitly** on genesis/G1 draft and mirror-build stages;

reserve low/medium for mechanical passes (lint fixes, copy swaps). Add an `effort` field to GAME-PLAN.md stage config. Re-audit the WebDesigner anti-slop prompt block for redundancy against 4.8 (the "harness diet" discipline from the prior report).

Source: <https://platform.claude.com/docs/en/build-with-claude/prompt-engineering/prompting-claude-opus-4-8> and <https://www.anthropic.com/news/claude-opus-4-8>

2. NEW — Fable 5 leads WebDev Arena (1653 Elo) over Opus 4.8 (1561); pick the model per stage.

What it is: Anthropic's Mythos-class Fable 5 became the WebDev Arena leader at 1653 Elo, ~92 points clear of Opus 4.8, per July 2026 rankings.

Why it matters: For raw single-shot UI generation quality, the strongest build model is no longer the default Opus. Model choice is now a mirror-score input.

How EdgePointe applies it: A/B the mirror-build agent on Fable 5 vs Opus 4.8 against our scored HC/Athens exemplars; if Fable wins on originality/design-quality, make it the builder and keep Opus for the skeptical evaluator (separation of concerns — the strongest generator shouldn't grade itself). Low risk to test, potentially a free quality bump across every build.

Source: <https://blog.logrocket.com/ai-dev-tool-power-rankings/> and <https://www.anthropic.com/news/claude-opus-4-8>

3. NEW/HIGH-VALUE — Impeccable: a deterministic (no-LLM) anti-slop detector + design vocabulary that inherits your tokens.

What it is: Paul Bakaus's Impeccable (44K+ GitHub stars) ships a `detect` CLI that runs **44 deterministic rules with no LLM** — catching gradient text, side-stripe borders, purple palettes, overused fonts, flat type hierarchy, WCAG contrast violations, nested cards, everything-centered layouts, bounce easing, and animating layout properties. It also gives the agent ~17–23 invocable commands (audit, critique, polish, simplify, animate) and "the missing design vocabulary," and explicitly *inherits your existing tokens/components* rather than inventing a parallel design system. Install: `npx impeccable install` then `/impeccable init`.

Why it matters: This is the deterministic pre-lint we sketched in the prior report (Finding 4's "detect Inter/Roboto in fontFaceCss"), already built, free, and far broader — and unlike an LLM evaluator it can't talk itself into a pass. The token-inheritance behavior means it won't fight Vantage's theme.json.

How EdgePointe applies it: Wire `impeccable detect` into the genesis lint gate as a **fast, cheap, deterministic pre-check that runs before the LLM mirror-score** — fail-fast on the 44 rules, only spend evaluator tokens on drafts that already pass. Map its rule IDs onto our existing L-16/L-17 taste lints; adopt its command vocabulary (`/polish`, `/simplify`) as named passes in clone-polish. Highest impact/effort ratio this cycle.

Source: <https://github.com/pbakaus/impeccable> and <https://github.com/pbakaus/impeccable/blob/main/README.md>

4. NEW — Dynamic Workflows: a native runtime for the mirror-score + parallel-builder loop.

What it is: Launched May 28, 2026, Dynamic Workflows let Claude write a JavaScript orchestration script that fans work across subagents (up to 16 concurrent, 1,000 total per run) in the background while the session stays responsive. The loop, branching, and intermediate results live in *script variables* — not the model's context — and the pattern natively supports "independent agents adversarially review each other's findings before they're reported" and "iterate until answers converge."

Why it matters: Our mirror-score harness (builder → skeptical evaluator → revise, 5–15 iterations) and the measure-first parallel-builder pattern are exactly this shape. Running them as a Dynamic Workflow keeps evaluator context clean (less drift), makes iteration-until-threshold a script `while` loop, and lets per-section builders fan out for free.

How EdgePointe applies it: Port `deep-evaluate.mjs` + the build loop into a Dynamic Workflow: script holds `while (score < target && i < max)`; parallel per-page/per-section builder subagents; the skeptical evaluator runs adversarial review in its own context each round. This is the "loop engineering" recommendation from the prior report, now with a supported orchestration layer instead of hand-rolled scripts.

Source: <https://claude.com/blog/introducing-dynamic-workflows-in-claude-code> and <https://code.claude.com/docs/en/workflows>

5. UPDATED — Official frontend-design skill: 277K → 565K installs; Space Grotesk now explicitly banned; "Elite 2026" variant.

What it is: Anthropic's frontend-design skill roughly doubled installs since the last run and now names Space Grotesk in its banned list (Inter, Roboto, Arial, system fonts, Space Grotesk). It prescribes distinctive display+body pairings (Playfair Display/Crimson Pro editorial; IBM Plex/Source Sans 3 technical; Bricolage Grotesque/Newsreader distinctive), weight extremes (100/200 vs 800/900), 3x+ size jumps, and "pick one distinctive font and use it decisively." An "Elite frontend-design-2026" variant (added April 18) carries more aggressive typography/layout rules for premium portfolios and landing pages.

Why it matters: Confirms the slop fingerprint keeps moving (Space Grotesk was a *recommended* alternative last cycle, now it's banned) — static blocklists rot. The Elite variant is a ready-made higher bar for our premium tenants.

How EdgePointe applies it: Refresh the WebDesigner banned-fonts list to include Space Grotesk; treat the blocklist as a maintained, dated artifact (re-pull the SKILL.md each report cycle). Evaluate the "Elite 2026" variant as the base taste layer for Concierge-tier builds (HC, Athens) over the standard skill.

Source: <https://github.com/anthropics/claude-code/blob/main/plugins/frontend-design/skills/frontend-design/SKILL.md> and <https://claudemarketplaces.com/skills/anthropics/skills/frontend-design>

6. UPDATED — Screenshot→critique loops: Playwright CLI now preferred over MCP (4x fewer tokens).

What it is: Microsoft now recommends the Playwright *CLI* over the MCP for coding agents because CLI uses ~4x fewer tokens per session; the "give Claude eyes" loop (change → navigate → screenshot → critique → fix → re-verify) is otherwise unchanged and remains table stakes, packaged as "Eyes"/visual-feedback skills.

Why it matters: A cheaper transport for the exact visual-QA loop we planned to make executable — lower token cost per QA sweep across many tenants adds up.

How EdgePointe applies it: Build the planned `visual-qa` skill (shots every page at 3 viewports, files a critique before any "done") on the Playwright **CLI** backend rather than MCP where feasible; keep Chrome DevTools MCP for the heavier computed-style/Lighthouse/a11y tier (unchanged from prior report).

Source: (Playwright MCP vs CLI) <https://testdino.com/blog/playwright-mcp> and <https://ap7i.com/posts/giving-claude-code-eyes-with-playwright-mcp/>

7. STABLE — Claude Design (web-capture prototyping) crossed 1M users first week; still the fastest G0 concept path.

What it is: Claude Design (Anthropic Labs, Opus 4.7 vision) — prompt-to-prototype with a web-capture tool that grabs elements from a live site so prototypes sit inside the real product UI; inline comments, direct text edits, adjustment knobs, hand-off to Claude Code. Reported 1M+ users in its first week; still research preview (Pro/Max/Team/Enterprise).

Why it matters: Unchanged recommendation from prior report, now with adoption proof — the cheapest way to show a client 2–3 on-brand directions before spending engine build time.

How EdgePointe applies it: Use for pre-genesis concept rounds (capture client's current site → 2–3 directions → owner inline comments → winner feeds genesis as the reference). A cheaper G0 than three staging drafts.

Source: <https://www.anthropic.com/news/claude-design-anthropic-labs> and <https://support.claude.com/en/articles/12138966-release-notes>

8. STABLE — The core doctrine holds: taste-context up front + a separate seeing evaluator; measure-first cloning; agent separation.

What it is: No reversal this cycle on the pillars established 07-16 — the GAN-style generator/evaluator harness (Anthropic Engineering, Mar 2026), measure-first cloning via `getComputedStyle` extraction, sprint contracts, negative prompting, file-based DESIGN.md context, and axe-core WCAG passes all remain current best practice and still map 1:1 onto our machinery.

Why it matters: Confirms the prior report's adopt-now list is still valid; nothing there was invalidated in two days.

How EdgePointe applies it: Continue executing the 07-16 top-5 (evaluator calibration/skepticism tuning, embed official aesthetics + banned-patterns, axe-core gate, per-page sprint contracts, pixel-diff baselines). See diff below.

Source: <https://www.anthropic.com/engineering/harness-design-long-running-apps>

Top 5 "adopt now" (impact vs effort)

- 1. Wire Impeccable's deterministic detect (44 rules) in as a fast pre-lint (Finding 3).** Free, no-LLM, can't self-approve, inherits our tokens; fails fast before we spend evaluator tokens. Highest impact/effort this cycle.
- 2. Codify model + effort selection per pipeline stage (Findings 1, 2).** Set explicit high effort on draft/build stages; A/B Fable 5 as the builder vs Opus 4.8; keep a separate model as skeptical evaluator. Near-zero effort, potential quality bump on every build.
- 3. Refresh the banned-fonts/patterns list (add Space Grotesk) and treat it as a dated, re-pulled artifact (Finding 5).** Copy-paste effort; the slop fingerprint moves monthly and stale blocklists silently rot.
- 4. Port the mirror-score loop onto Dynamic Workflows (Finding 4).** Moderate effort; native runtime for iterate-until-threshold + adversarial cross-review + parallel builders, with cleaner evaluator context.

5. **Build the `visual-qa` skill on the Playwright CLI backend (Finding 6).** Moderate effort; makes "screenshot+eyeball before done" executable and cheaper per sweep across the whole tenant book.

Diff vs prior report (2026-07-16)

- **New since last run:** Impeccable deterministic detector (F3); Dynamic Workflows as a loop runtime (F4); Opus 4.8 lower prompting need + strict `effort` (F1); Fable 5 leading WebDev Arena (F2); Playwright CLI > MCP on tokens (F6); frontend-design installs 277K→565K with Space Grotesk newly banned + Elite-2026 variant (F5).
- **Still valid, carry forward:** all 18 prior findings hold. The prior top-5 (evaluator calibration, embed aesthetics+banned-patterns, axe-core WCAG gate, sprint contracts, pixel-diff baselines) is unchanged and should keep executing.
- **Net shift in guidance:** last cycle was "close the loop with a seeing evaluator + steer with taste-context." This cycle adds a layer *below* the LLM evaluator (deterministic detection via Impeccable) and *around* it (Dynamic Workflows runtime + explicit model/effort selection). The evaluator is no longer the only quality gate — cheaper deterministic checks now front-run it.

Sources list

- <https://platform.claude.com/docs/en/build-with-claude/prompt-engineering/prompting-claude-opus-4-8> — Anthropic: Prompting Claude Opus 4.8 (less design prompting; strict `effort`)
- <https://www.anthropic.com/news/claude-opus-4-8> — Anthropic: Introducing Claude Opus 4.8
- <https://blog.logrocket.com/ai-dev-tool-power-rankings/> — LogRocket: AI dev tool power rankings (WebDev Arena Elo, Jul 2026)
- <https://github.com/pbakaus/impeccable> — Impeccable (Paul Bakaus): deterministic anti-slop rules + design vocabulary
- <https://github.com/pbakaus/impeccable/blob/main/README.md> — Impeccable README (44 detect rules, install)
- <https://claude.com/blog/introducing-dynamic-workflows-in-claude-code> — Anthropic: Introducing Dynamic Workflows
- <https://code.claude.com/docs/en/workflows> — Claude Code docs: orchestrate subagents with dynamic workflows
- <https://github.com/anthropics/claude-code/blob/main/plugins/frontend-design/skills/frontend-design/SKILL.md> — Anthropic frontend-design skill (banned fonts incl. Space Grotesk)
- <https://claudemarketplaces.com/skills/anthropics/skills/frontend-design> — frontend-design skill listing (installs, Elite-2026 variant)
- <https://www.anthropic.com/engineering/harness-design-long-running-apps> — Anthropic Engineering: generator/evaluator harness (still current doctrine)
- <https://www.anthropic.com/news/claude-design-anthropic-labs> — Anthropic: Claude Design (web-capture prototyping)
- <https://support.claude.com/en/articles/12138966-release-notes> — Claude release notes (Jul 2026)

- <https://ap7i.com/posts/giving-claude-code-eyes-with-playwright-mcp/> — Giving Claude Code eyes with Playwright MCP
- <https://testdino.com/blog/playwright-mcp> — Playwright MCP vs CLI token cost (2026)
- <https://platform.claude.com/cookbook/coding-prompting-for-frontend-aesthetics> — Anthropic Cookbook: frontend aesthetics prompt (baseline, carried forward)