



INTUITIONLABS / RESEARCH & PRACTICAL GUIDANCE

Anthropic Adaptyv Protein Design Competition Guide

Build practical AI for pharma and biotech with IntuitionLabs. We help life-science teams turn complex information and workflows into useful software, governed knowledge systems and AI tools.

Published by IntuitionLabs · 2026-09-19 · anthropic adaptyv protein design competition · proteinbase competition · claude protein design · protein engineering · wet-lab validation · ai drug discovery · open data · biotechnology competitions

Original article: <https://intuitionlabs.ai/articles/anthropic-adaptyv-protein-design-competition-guide>



In this article

1. Executive Summary
 2. Introduction and Background
 3. Competition Definition, Calendar, and Offer
 4. Track Decision Framework
 5. Eligibility, Support, and Weekly Operating Plan
 6. Validation and Selection
 7. IP, Licensing, Confidentiality, and Publication
 8. Implementation Guidance and Submission Package
 9. Data Analysis and Evidence
 10. Implications and Future Directions
 11. Frequently Asked Questions (FAQs)
 12. Conclusion
-

Executive Summary

The **Anthropic Adapyv Protein Design Competition** is a five-challenge, open-results program scheduled from **September 28 through October 31, 2026**. Applications for Anthropic-supported Tracks 1 and 2 close **September 24**, while the self-supported Track 3 begins with the competition. The terms describe an expected schedule, subject to change by the Sponsors: final design submission deadline **October 31**, experimental validation **November 30**, and results published **December 15** (proteinbase.com) (proteinbase.com). As of **September 19**, the targets, exact weekly times, sequence constraints, upload templates, and detailed selection workflow were not public. The organizers say requirements will appear when each problem opens (proteinbase.com).

The advertised **\$1 million in Claude credits** and **\$1 million in experimental validation** are aggregate program support, not cash prizes or per-team awards. The terms explicitly state that there is no cash prize (proteinbase.com). Track 1 is the strongest fit for experienced labs or companies that can sustain five weekly sprints: no more than **20 teams** may receive up to **\$50,000** in Claude credits for academic teams or **\$25,000** for industry teams (proteinbase.com). Track 2 suits one to three affiliated researchers seeking complimentary **Claude Max 20x** access, but support is selective. Track 3 permits eligible participants to use their own tools and compute. Only Track 1 has a stated reserved testing allocation, up to approximately **18 designs per challenge**; Tracks 2 and 3 have no guaranteed allocation (proteinbase.com).

The central decision is therefore not whether free validation sounds attractive. It is whether a team can accept **open publication of every tested design**, including sequences, predicted structures, methods, experimental measurements, and negative results (proteinbase.com) (journals.plos.org). Participants retain ownership, but grant a non-exclusive, worldwide, royalty-free, perpetual, irrevocable publication license (proteinbase.com). The terms also warn that publication may affect patentability (wipo.int). Teams with patent-sensitive sequences should complete invention review and any intended filings before submission, with jurisdiction-specific counsel.

The practical recommendation is to **enter Track 1** only if the team has mature protein-design operations, institutional authorization, and a pre-cleared IP plan; **enter Track 2** if a small affiliated team values Claude access and accepts uncertain testing; **enter Track 3** if it can self-fund compute or wants method freedom; and **defer** if public disclosure, weekly staffing, or unknown target requirements create unacceptable risk. Wet-lab results will be conditional on a sponsor-selected subset, not a random sample. Historical campaigns show why that matters: one Proteinbase challenge tested **1,196 designs**, of which **1,028 expressed** and **111 bound** (proteinbase.com), while peer-reviewed campaigns generated **15,000 to 100,000 candidates per target** before selection (nature.com). A credible submission should preserve the full generated-to-screened funnel, software versions, parameters, seeds, prompts, human decisions, and rejection reasons.

Introduction and Background

This report is an **entry, experiment-design, and intellectual-property decision guide** for protein-engineering leads, computational biologists, biotechnology research directors, technology-transfer counsel, and academic laboratory heads. It addresses the immediate question before **September 24, 2026**: whether to apply for sponsor support, which track to select, and whether the open publication model fits the team's scientific and commercial strategy.

The competition is not the same as deploying Anthropic's open optimization repository. Anthropic describes the event as five difficult problems, including species cross-reactivity, pH sensitivity, peptide-major histocompatibility complex specificity, and difficult target classes (anthropic.com). Its separate code release covers optimized biomolecular-modeling workloads, but the repository is expressly **not maintained** and does not accept contributions (github.com). Entrants may use Claude or other tools in Track 3; Claude use becomes a condition only when accepting Anthropic support.

From an adjacent-advisor perspective, the decision should be governed like a short, high-intensity research program, not treated as a software trial. IntuitionLabs describes its role as connecting artificial intelligence to authoritative sources with permissions, retrieval, citations, evaluation, and accountable operation (intuitionlabs.ai), and separately emphasizes measuring quality, reliability, risk signals, and support burden before scaling (intuitionlabs.ai). Applied here, that means documenting what was generated, why a human approved it, what the selectors tested, and which conclusions the resulting data can and cannot support. IntuitionLabs is not an entrant track, competition vendor, or comparison-table option.

Competition Definition, Calendar, and Offer

The program combines **compute support, a sequence-design competition, sponsor-directed candidate selection, automated synthesis and testing, and open data publication**. It is free to enter, but participation has real internal costs: scientist time, compute, institutional approvals, invention triage, method documentation, and post-result analysis.

What the program offers

- **Aggregate Claude support:** Anthropic is offering up to **\$1 million** in Claude credits across the program (anthropic.com).

- **Aggregate validation:** The sponsors describe wet-lab validation for more than **5,000 designs** across the program, not for each participant (anthropic.com).
- **Additional compute:** Anthropic's announcement also identifies up to **\$250,000 in Modal compute credits** as a program contribution (anthropic.com).
- **Prospective problems:** The problems are released sequentially, so entrants cannot fully optimize a target-specific pipeline before the event. This resembles blind-challenge controls such as CASP17, which requires experimental structures to remain unreleased until the applicable prediction period closes (predictioncenter.org).
- **Open tested results:** Positive and negative tested outcomes are intended for publication. That can add scientific value because unpublished null studies waste resources and slow cumulative learning (journals.plos.org).

Date-stamped uncertainty

As of **September 19, 2026**, target requirements and exact opening and closing times remained pending. The sponsors may jointly cancel, terminate, modify, or suspend a track, problem, target, or timeline (proteinbase.com). Staffing and compute reservations should therefore be flexible, with a change-control owner assigned.

Track Decision Framework

The three tracks differ primarily in **who may enter, who may receive Anthropic support, and whether any wet-lab capacity is reserved**. They do not represent three scientific difficulty levels.

Table 1 compares the decision-relevant terms. "Considered" is deliberately different from "guaranteed."

Decision factor	Track 1: labs and companies	Track 2: individuals and small teams	Track 3: open track
Best fit	Experienced protein-design lab or company with relevant expertise or publications	One affiliated researcher or a team of up to three	Eligible participant outside Tracks 1 and 2
Affiliation	Institutional or company email; an authorized support letter may be requested	Active institution or company affiliation and official email	No Track 1 or 2 status; general eligibility still applies
Claude support	Up to \$50,000 academic or \$25,000 industry	Complimentary Claude Max 20x for the competition	Self-supported; Claude is optional
Selection limit	Maximum 20 teams (proteinbase.com)	Support recipients selected by Anthropic	No support-recipient selection
Wet-lab position	Public page states up to approximately 18 designs per challenge reserved	Considered through the joint workflow, no guaranteed allocation	Considered through the joint workflow, no guaranteed allocation

Decision factor	Track 1: labs and companies	Track 2: individuals and small teams	Track 3: open track
Tool condition	Claude use required if support is accepted	Claude use required if support is accepted	Any design tools may be used
Publication consequence	Tested designs and results become public under the program terms	Same	Same
Primary tradeoff	Highest support and testing certainty, highest weekly delivery burden	Low direct support cost, uncertain validation	Maximum tool freedom, self-funded compute and uncertain validation

The table shows why **Track 1 is not automatically the best choice**. A reserved allocation has value only if the team can produce and review credible candidates in all five windows. The public allocation implies up to about **90 reserved designs per selected team** if a team uses approximately 18 slots in each of five challenges, but this is an arithmetic ceiling, not a promise. The binding terms preserve sponsor discretion over testing quantities. Track 1 applicants should therefore budget around the downside case as well as the headline allocation.

An enter or defer decision tree

- **Choose Track 1** if the organization has significant protein-design evidence, can secure an authorized applicant and institutional backing, can staff five consecutive sprints, and has cleared publication and patent issues.
- **Choose Track 2** if a small affiliated team can benefit from Claude Max access, has a reproducible workflow, and views wet-lab selection as upside rather than an entitlement.
- **Choose Track 3** if the participant is eligible, wants to use non-Claude methods, or does not need sponsor compute. It remains subject to open publication for tested designs.
- **Defer** if the candidate sequences must remain confidential, foreign patent options are not yet protected, staff cannot cover each weekly review, or the project needs guaranteed testing.
- **Apply conservatively** if institutional authority is uncertain. Support selection considers expertise, publication record, and scientific novelty, and Track 1 may require an institutional support letter (proteinbase.com).

Eligibility, Support, and Weekly Operating Plan

General eligibility starts with **age, residence or domicile, sanctions-related restrictions, and acceptance of the terms**. Individuals must be at least 18 or the higher local age of majority. Track 1 and Track 2 require proof of affiliation through an official institutional or company email (proteinbase.com). Support applicants must also be in a country or region where Claude is available for permissible use; Anthropic maintains current access lists separately (anthropic.com).

The competition terms exclude legal residents and domiciled entities in **Belarus, China, Cuba, Iran, Myanmar, North Korea, Russia, Sudan, Syria, Crimea, Donetsk, and Luhansk**. Teams should screen the exact participant and entity against the terms rather than infer eligibility from a generic country list. The U.S. Office of Foreign Assets Control notes that it does not maintain one universal country blacklist because restrictions vary by program (ofac.treasury.gov). A possible name match on the U.S. Consolidated Screening List also requires additional due diligence (trade.gov).

Table 2 converts the published dates into a staffing plan. The exact intraday cutoffs remained unpublished on September 19.

Window	Published milestone	Internal operating target	Minimum accountable roles
Sep 16 to Sep 24	Support applications	Complete track, affiliation, support letter, IP, and compute approvals 24 hours early	Principal investigator, institutional official, IP counsel, compute owner
Before Sep 28	Selected support recipients notified	Confirm account access and contingency compute immediately	Team lead, platform administrator
Sep 28 to Oct 4	Challenge 1	Dry-run provenance, screening, ranking, and reviewer sign-off	Designer, reviewer, data steward
Oct 5 to Oct 11	Challenge 2	Reuse validated pipeline, log target-specific changes	Same core team plus target expert
Oct 12 to Oct 18	Challenge 3	Audit selection funnel and compute consumption	Team lead, compute owner
Oct 19 to Oct 25	Challenge 4	Preserve rejected candidates and reasons	Data steward, independent reviewer
Oct 26 to Oct 31	Challenge 5 and final deadline	Freeze manifest, verify submission receipt, archive environment	All accountable owners
By Nov 30	Target validation date	Predefine analysis before seeing outcomes	Assay scientist, statistician
Dec 15	Target publication date	Reconcile public data, invention records, and report	Data steward, counsel, scientific lead

This calendar leaves no recovery week. Each problem is revealed one at a time, each window lasts roughly one week, and one window closes as the next begins. A practical team should therefore pre-build container images, sequence-quality checks, ranking notebooks, provenance templates, and reviewer forms before September 28. It should also designate backups for the scientific approver and Proteinbase submitter.

Support and account diligence

- **Confirm the recipient account terms.** The competition terms say Claude inputs and outputs are governed by the terms applicable to the recipient account.

- **Do not assume one retention policy.** Anthropic says standard application programming interface inputs and outputs are normally deleted within **30 days**, subject to exceptions, but that may not govern consumer Claude Max accounts (privacy.claude.com).
- **Document human review.** Anthropic's commercial terms require customers to assess where human review is appropriate before using or sharing outputs (anthropic.com).
- **Respond promptly.** A selected applicant can be replaced after failing to respond within **five business days** of the first notification attempt.

Validation and Selection

The experimental offer is substantial, but the results will answer a narrower question than many entrants may assume. The organizers will select designs with a **Claude-based workflow**, will not rely on a single in silico metric, and plan to disclose the exact workflow only after the competition. Every design must be reviewed by the submitting researcher before submission. Selection therefore occurs at least twice: first inside the entrant's pipeline, then inside the sponsors' workflow. NIST recommends that human-oversight processes be defined, assessed, and documented under organizational policies (airc.nist.gov), reinforcing the need for substantive review.

The correct funnel

Teams should report five denominators rather than a single "hit rate":

1. **Designed:** all sequences emitted by the generative process.
2. **Screened:** sequences passing safety, quality, duplication, and basic feasibility checks.
3. **Submitted:** the entrant-approved subset delivered to Proteinbase.
4. **Selected and expressed:** sponsor-selected sequences that synthesize and yield measurable protein.
5. **Experimentally positive:** expressed molecules meeting the disclosed assay criterion.

This separation matters because selected wet-lab candidates are not a random sample. A prior Adaptyv and muni hackathon received **141 designs**, ranked them computationally, and tested the top **100**; **37** were reported as binders (proteinbase.com) (proteinbase.com). That **37%** is a rate within the selected test set, not within every sequence generated.

Peer-reviewed evidence shows the same selection compression. One prospective study designed **15,000 to 100,000 binders for each of 13 sites**, prioritized candidates using folding and interface metrics, and in some campaigns obtained fewer than **10 binders from 100,000 designs** (nature.com) (nature.com). Another study's retrospective training data included about **1 million experimentally characterized designs across 10 targets**, with **15,000 to 100,000** tested per target (nature.com) (nature.com). The same paper explains that only a small fraction of computational designs typically bind strongly enough for experimental detection (nature.com).

What can and cannot be inferred

- **Can infer:** how the tested, sponsor-selected subset performed in the named assays under the disclosed conditions.
- **Cannot infer:** an unbiased success rate for every generated sequence unless the full generation and selection denominators are reported.
- **Can compare cautiously:** methods evaluated on the same targets, assay definitions, budgets, and selection pipeline.
- **Cannot compare directly:** percentages from different targets or campaigns without accounting for target difficulty, candidate counts, filtering, expression, and success thresholds.
- **Can learn from negatives:** published failures can improve future filtering and reduce repeated dead ends.
- **Cannot treat in silico rank as biological truth:** PDBench notes that sequence-recovery accuracy does not capture real-world method utility and that a static single structure does not represent protein behavior in solution (academic.oup.com) (academic.oup.com).

Blind-challenge practice provides a useful analogy. EMBL-EBI describes CAPRI as a blind test of protein-protein docking algorithms and warns that using both bound component structures biases docking too strongly toward the correct answer (ebi.ac.uk) (ebi.ac.uk). CASP17 likewise requires experimental structures to remain unreleased until the prediction window closes (predictioncenter.org). The competition's sequential unrevealed targets improve prospective relevance, but post-competition interpretation still depends on the unpublished selection workflow.

IP, Licensing, Confidentiality, and Publication

The core legal distinction is **ownership versus permission**. Participants retain ownership of their designs, yet grant the sponsors a non-exclusive, worldwide, royalty-free, perpetual, and irrevocable license for publication purposes (proteinbase.com). Tested materials are not participant-confidential under the publication terms. Ownership therefore does not preserve secrecy.

Table 3 is an operational checklist, not legal advice. Teams should apply it with their institution's technology-transfer office or qualified counsel.

Issue	Competition consequence	Action before submission
Participant ownership	Ownership remains with the participant	Confirm employee, student, collaborator, sponsor, and background-IP obligations
Sponsor license	Broad, perpetual publication permission for tested materials	Ensure every contributor has authorized the grant
Published content	Sequence, predicted structure, method, measurements, positive and negative results	Remove third-party confidential inputs and identify background assets

Issue	Competition consequence	Action before submission
Patent timing	Terms warn publication may affect patentability	Triage inventions and make intended filings before submission
ODC-By 1.0	Database reuse is allowed with attribution obligations	Record dataset title, source URL, license URI, version, and access date
Confidential sponsor information	Prelaunch targets and assay details may be restricted to competition use	Limit access, label records, and avoid external reuse
Model-account terms	Input, output, retention, and training terms depend on the account	Record account type and applicable terms at each challenge
Biosecurity	Every design sent for synthesis is screened	Run internal screening first and retain reviewer sign-off

The checklist shows why filing strategy must precede the first submission. The U.S. Patent and Trademark Office says a U.S. provisional application may be filed up to **12 months after an inventor's public disclosure**, but also warns that the same disclosure may prevent foreign patenting ([uspto.gov](https://www.uspto.gov)) ([uspto.gov](https://www.uspto.gov)). The European Patent Convention defines the state of the art broadly as material made public by written or oral description, use, or another way before filing ([epo.org](https://www.epo.org)). EPO guidance treats internet material as public from the posting date ([epo.org](https://www.epo.org)). Canada's patent office similarly advises not disclosing an invention publicly before filing ([ised-isde.canada.ca](https://www.ised-isde.canada.ca)). These differences make a generic "12-month grace period" assumption unsafe.

What ODC-By does and does not do

The **Open Data Commons Attribution License 1.0** governs database rights, not necessarily the rights in each individual content item (opendatacommons.org). It permits commercial use and does not exclude a field of endeavor, but it does not itself license patents over the database or contents (opendatacommons.org) (opendatacommons.org). Public conveyance of the database or a derivative requires the license or its URI to accompany it (opendatacommons.org). A downstream user should therefore track database attribution, content rights, and patent rights as separate layers.

Implementation Guidance and Submission Package

The official materials did not yet publish per-challenge file formats, sequence limits, ranking fields, or an upload template as of September 19. The current route required a Proteinbase account, and the page asks participants to describe methods and design choices in detail. Teams should prepare a richer internal package than the minimum portal fields so that omissions in a one-week sprint do not destroy reproducibility.

Recommended evidence package

- **Identity and authority:** legal participant name, track, institutional affiliation, official email, authorized submitter, and support letter if requested.

- **Target snapshot:** target name, sequence or structure identifier, challenge requirement version, release timestamp, and local checksum. NSF biological-sciences guidance treats provenance, archiving, access timing, standards, metadata, and accountability as explicit data-plan categories ([nsf.gov](https://www.nsf.gov)).
- **Method manifest:** model names, exact versions, repository commits, container digest, dependencies, inference settings, templates, and scoring functions.
- **Randomness:** seeds, sampling temperatures, stochastic filters, rerun policy, and candidate deduplication rules. A computational-biology benchmarking guide specifically recommends reporting parameter values, random seeds, and software versions (link.springer.com).
- **Full funnel:** generated, safety-screened, technically valid, ranked, researcher-reviewed, submitted, sponsor-selected, expressed, and positive counts.
- **Human review:** reviewer identity, competence, date, decision, rationale, escalations, and rejected-design reasons. NIST recommends that human-oversight processes be defined, assessed, and documented under organizational policies (airc.nist.gov). It also calls for model output to be interpreted within its intended context (airc.nist.gov).
- **Biosecurity record:** internal sequence-screening result, flagged similarities, disposition, and authorized approval. The competition's sponsor screening before DNA synthesis is a second control, not a substitute for the submitting researcher's duty. NIH procurement guidance permits framework adherence to be demonstrated through written or public provider attestation (grants.nih.gov).
- **IP disposition:** invention disclosure number, ownership check, filing decision, filing date if applicable, and approved public fields.
- **Compute ledger:** model calls, graphics-processing-unit hours, token usage, direct cost, sponsor credit, failed runs, and peak capacity.
- **Submission receipt:** final sequence list, rank order, hashes, portal timestamp, confirmation, and any corrected resubmission.
- **Analysis plan:** assay endpoint, success threshold, handling of non-expression and missing data, denominator definitions, and comparisons to be made after results.

Reproducibility and provenance

Provenance means recording the entities, activities, and people involved in producing a data item (w3.org). The National Academies recommends preserving the data, study methods, and computational environment required to repeat an analysis (ncbi.nlm.nih.gov). NIST's reproducibility guidance likewise recommends an executable script that exactly reproduces reported results (tsapps.nist.gov).

Dataset versions also matter after publication. DataCite recommends assigning and linking a new digital object identifier when a research resource changes substantially (support.datacite.org). Creative Commons recommends six human-readable and machine-readable metadata values for scientific-data licensing and attribution (creativecommons.org). Proteinbase exposes sequence, design-method, and evaluation records, so teams should capture the publication version and access date before running post hoc analyses.

Data Analysis and Evidence

No defensible competition-wide success rate exists before the targets, selection workflow, assays, and results are published. An expected-value model should therefore use team-specific probabilities and explicit uncertainty rather than invented benchmark odds.

Quantitative evidence from prior work

- **Nipah competition funnel: 1,196 tested, 1,028 expressed, and 111 bound**, corresponding to **86% expression** and **9.3% binding among tested designs**. The binding rate among expressed proteins is approximately **10.8%**, calculated as 111 divided by 1,028. It is not the rate among all generated candidates.
- **Large-scale 2017 screen**: researchers designed and tested **22,660 miniproteins** targeting two proteins and identified **2,618 high-affinity binders** ([nature.com](#)) ([nature.com](#)). That is about **11.6%** across the reported tested set, not a transferable forecast for the 2026 targets.
- **RFdiffusion campaign**: one study selected **95 designs per target** for experimental characterization after computational generation and filtering ([nature.com](#)). The authors reported that improved filtering raised experimental success rates by about **two orders of magnitude** relative to prior work ([nature.com](#)).
- **BindCraft range**: a 2025 paper reported experimental success rates from **10% to 100%**, illustrating target and campaign heterogeneity rather than a universal expectation ([nature.com](#)).
- **Computational efficiency**: ProteinMPNN generated a minibinder sequence in about **2 CPU-seconds**, versus about **350 CPU-seconds** for Rosetta design in the reported workflow ([nature.com](#)). Faster generation can increase the designed denominator without increasing wet-lab capacity, making selection policy more consequential.

Expected-value worksheet

For each track, estimate:

Expected value = probability of receiving support × usable support value + probability of wet-lab selection × usable testing value + learning value + publication value - staff cost - uncovered compute cost - IP and disclosure cost - opportunity cost.

Use internal inputs, not list-price fiction:

- **Support value**: count only credits the team can consume within the competition and only for eligible work.
- **Testing value**: use the team's avoided marginal cost for equivalent assays, adjusted for the probability that its designs are selected and express successfully.
- **Staff cost**: include five weekly design, review, biosecurity, submission, and documentation sprints.

- **IP cost:** include invention harvesting, outside counsel where needed, filing fees, and the strategic cost of public negative and positive data.
- **Learning value:** value new labeled data only if the assay and metadata can change the next design decision.
- **Publication value:** distinguish citable open data from peer-reviewed authorship, which the terms do not guarantee.
- **Risk adjustment:** model schedule changes, support non-selection, zero sponsor-selected designs in Tracks 2 or 3, non-expression, and target mismatch.

Run at least three scenarios: **downside**, with no support or testing; **base**, using the team's conservative selection probability; and **upside**, using only the stated allocation ceilings. Do not use the aggregate **\$2 million** support headline as participant value. It cannot be booked by an individual team.

Implications and Future Directions

The competition could produce a useful open dataset because it combines previously unseen targets, human-reviewed submissions, automated testing, and publication of negative outcomes. A scientific consensus warns that unpublished null results waste resources and slow research (journals.plos.org). Scientific value will depend on complete funnel denominators, assay definitions, target metadata, selection rationales, and versioned methods. NIH's data-management policy treats data needed to validate and replicate findings as scientific data (grants.nih.gov).

The official competition page states that the exact Claude-based selection workflow will be made public after the competition.

Biosecurity and human review will also shape future norms. Every design sent for synthesis is screened before DNA synthesis, while current U.S. policy calls for comprehensive, scalable, verifiable synthetic nucleic-acid procurement screening (whitehouse.gov). NIST's governance framework separately calls for documented human-oversight processes (airc.nist.gov). A strong program record should show both machine screening and accountable researcher judgment, not merely a portal acceptance.

Finally, open data will be most useful when it is citable and versioned. DataCite, Crossref, and STM recommend citing datasets with persistent identifiers (datacite.org). DataCite also recommends a linked new identifier for a substantially changed resource (support.datacite.org). Participants should plan a post-publication record that connects the public dataset, internal run manifest, code revision, and any subsequent paper without implying that the database license grants patent rights.

Frequently Asked Questions (FAQs)

How does a team enter the Anthropic Adaptevy competition?

Tracks 1 and 2 require the joint application by **September 24, 2026**, an official institutional or company email, acceptance of the competition terms, and acceptance of the Anthropic support terms if support is requested. Track 1 may require an authorized institutional support letter. Designs are later submitted through Proteinbase

under the requirements published for each problem. Track 3 is self-supported and begins with the competition.

Are international participants eligible?

Eligibility is global subject to the stated age, residence, domicile, entity, sanctions, export-control, and Claude-availability restrictions. The competition's excluded-jurisdiction list controls, not a general assumption about international access. A team should also verify that every named participant and entity is eligible.

Must an entrant use Claude?

No. Claude is optional for general entry and Track 3 may use other tools. Claude use is a condition of receiving Anthropic support in Tracks 1 and 2.

Is wet-lab testing guaranteed?

Only Track 1 has a stated reserved allocation, up to approximately 18 designs per challenge for selected teams. Tracks 2 and 3 have no guaranteed allocation. Even for Track 1, teams should read the binding terms and plan for sponsor discretion over testing quantities.

What validation does Adaptyv perform?

The program promises synthesis and experimental characterization of sponsor-selected designs, but the exact targets, assay requirements, and per-challenge specifications were not public on September 19. **Experimental validation is not synonymous with clinical relevance**, therapeutic efficacy, or an unbiased benchmark of all generated sequences.

Who selects designs for testing?

Anthropic and Adaptyv will use a jointly agreed Claude-based workflow. Selection will not rely on one in silico metric, and the exact workflow is scheduled for disclosure after the competition. Entrants must separately review every design before submission.

Who owns the designs?

Participants retain ownership. For tested designs, however, the sponsors receive broad publication permission and intend to release the sequence, predicted structure, method, experimental measurements, and negative as well as positive results under ODC-By. Ownership should not be confused with confidentiality.

Can a participant still seek patents?

The terms preserve the participant's ability to patent, license, or commercialize, but warn that publication may affect patentability and place responsibility for filings on the participant. WIPO advises filing before public disclosure because disclosure generally becomes prior art ([wipo.int](https://www.wipo.int)). Because grace periods and novelty rules differ by jurisdiction, patent-sensitive teams should decide and act before submission ([uspto.gov](https://www.uspto.gov)).

What should be preserved for reproducibility?

Preserve target inputs, exact model and software versions, code commits, environment files, parameters, random seeds, prompts, full generated and filtered candidate sets, human-review decisions, compute usage, submitted sequence hashes, sponsor-selection status, expression outcomes, assay definitions, and final measurements. NIH defines scientific data to include data needed to validate and replicate research findings (grants.nih.gov). Reproducibility guidance recommends preserving the data, methods, and computational environment (ncbi.nlm.nih.gov) and reporting random seeds and software versions (link.springer.com).

Conclusion

The Anthropic Adaptyv competition offers an unusual combination of sponsor compute, prospective targets, automated wet-lab capacity, and open positive and negative results. Its value is largest for teams that already have a robust design pipeline but lack comparable experimental throughput. The event is not a cash-prize contest, and the aggregate support figures should not be treated as per-team economic value.

The track choice is straightforward when framed around constraints. Track 1 fits experienced labs and companies that can meet five weekly deadlines and accept public disclosure. Track 2 fits small affiliated teams that value Claude access but can tolerate uncertain testing. Track 3 fits self-supported participants that want tool freedom. Teams should defer when confidentiality, patent timing, staffing, compute readiness, or the absence of guaranteed testing outweighs learning and validation value.

The strongest entry will combine scientific ambition with disciplined governance: a pre-submission IP decision, clear institutional authority, an auditable generated-to-tested funnel, independent researcher review, biosecurity screening, fixed denominator definitions, and a reproducible environment. Until targets and the selection workflow are disclosed, no team should invent a success forecast. The defensible approach is to enter with explicit downside, base, and upside scenarios, then interpret the published results as outcomes for a selected experimental subset rather than as a universal measure of protein-design capability.

Source links

Links cited in this article, in order of first appearance. Full addresses are provided for readers using extracted text.

1. <https://proteinbase.com/competitions/anthropic-adaptyv-2026/terms>
2. <https://proteinbase.com/competitions/anthropic-adaptyv-2026>
3. <https://journals.plos.org/plosbiology/article?id=10.1371%2Fjournal.pbio.3003368>
4. https://www.wipo.int/en/web/patents/faq_patents
5. <https://proteinbase.com/collections/nipah-binder-competition-results>
6. <https://www.nature.com/articles/s41586-022-04654-9>
7. <https://www.anthropic.com/research/claude-uplifts-biomolecular-modeling>
8. <https://github.com/anthropics/uplifting-biomolecular-modeling>
9. <https://intuitionlabs.ai/>
10. <https://intuitionlabs.ai/articles/modern-biotech-lab-automation-ai>
11. https://predictioncenter.org/casp17/targets_submission.cgi

12. <https://www.anthropic.com/supported-countries>
13. <https://ofac.treasury.gov/sanctions-programs-and-country-information/where-is-ofacs-country-list-what-countries-do-i-need-to-worry-about-in-terms-of-us-sanctions>
14. <https://www.trade.gov/consolidated-screening-list>
15. <https://privacy.claude.com/en/articles/7996866-how-long-do-you-store-my-organization-s-data>
16. <https://www.anthropic.com/legal/commercial-terms>
17. <https://airc.nist.gov/airmf-resources/airmf/5-sec-core/>
18. <https://proteinbase.com/collections/adaptyv-x-muni-hackathon-ai-agents-vs-humans>
19. <https://www.nature.com/articles/s41467-023-38328-5>
20. <https://academic.oup.com/bioinformatics/article/39/1/btad027/6986968>
21. <https://www.ebi.ac.uk/pdbe/complex-pred/capri/call-for-targets/>
22. <https://www.uspto.gov/patents/basics/apply/provisional-application>
23. <https://www.epo.org/en/legal/epc/2020/a54.html>
24. https://www.epo.org/en/legal/guidelines-pct/2025/g_iv_6_4.html
25. <https://ised-isde.canada.ca/site/canadian-intellectual-property-office/en/ip-roadmap-your-path-getting-patent-grant>
26. <https://opendatacommons.org/licenses/by/1-0/>
27. <https://www.nsf.gov/bio/data-management-plans>
28. <https://link.springer.com/article/10.1186/s13059-019-1738-8>
29. <https://www.grants.nih.gov/grants/guide/notice-files/NOT-OD-25-012.html>
30. <https://www.w3.org/TR/prov-overview/>
31. <https://www.ncbi.nlm.nih.gov/books/NBK547531/>
32. https://tsapps.nist.gov/publication/get_pdf.cfm?pub_id=936953
33. <https://support.datacite.org/docs/versioning>
34. <https://creativecommons.org/2026/04/20/licensing-best-practices-for-the-sharing-of-scientific-data/>
35. <https://www.nature.com/articles/nature23912>
36. <https://www.nature.com/articles/s41586-023-06415-8>
37. <https://www.nature.com/articles/s41586-025-09429-6>
38. <https://www.grants.nih.gov/policy-and-compliance/policy-topics/sharing-policies/dms/policy-overview>
39. <https://www.whitehouse.gov/presidential-actions/2025/05/improving-the-safety-and-security-of-biological-research/?query=11-page=72>
40. <https://datacite.org/blog/joint-statement-on-research-data/>
41. <https://intuitionlabs.ai/articles/preclinical-vs-clinical-research>

THE PUBLISHER

About IntuitionLabs

IntuitionLabs is an AI consulting, custom software development and data engineering firm serving pharmaceutical, biotechnology, medical-device and other life-science organizations. We work with clinical, regulatory, medical-affairs, commercial, quality and IT teams to connect technology decisions with the work people need to accomplish.

AI consulting and adoption

Our AI enablement services cover readiness assessments, use-case selection, governance and policies, team workshops, adoption measurement and ongoing advisory support. We help organizations structure the information layer behind AI: source material, context, permissions and maintained knowledge that make generated answers useful and reviewable. Private LLM inference and hosted AI options support teams evaluating how to operate AI with appropriate control over their data and infrastructure.

Software, data and life-science workflows

IntuitionLabs develops custom software for pharma and biotech, integrates enterprise systems, and builds data engineering and business intelligence solutions. Areas of focus include AI agents, regulatory research, medical writing, medical affairs, CMC information, competitive intelligence and clinical-document workflows. Our eTMF intelligence work includes cross-system reconciliation and inspection-readiness support.

Enterprise platforms and regulated delivery

We provide Veeva services, application support, managed services, integrations and custom applications, alongside enterprise content work involving platforms such as Egnyte. For regulated workflows, our services include GxP enablement, computer-system validation and software development addressing 21 CFR Part 11 requirements. The applicable controls, validation responsibilities and acceptance criteria are defined for each engagement.

Work with IntuitionLabs

Explore [AI enablement](#), [pharma and biotech software development](#), [data engineering and BI](#), and [Veeva services](#). [Contact IntuitionLabs](#) to discuss your workflow, information sources and implementation needs.

IntuitionLabs publishes educational research to help life-science teams make informed technology decisions. Coverage of a product or organization does not imply a client relationship, endorsement or partnership.

Website: <https://intuitionlabs.ai>

This article is educational. Verify consequential medical, legal, financial, regulatory and technical decisions with appropriate professionals and the cited primary sources. Company services are described separately from the article's research findings.