

Designing University Startup Accelerators: Linking Program Architecture to Documented Venture Evidence

Priyanshu Ghosh, UW-Madison

Priyanshu Ghosh is a researcher and entrepreneur focused on the intersection of data science, innovation systems, and venture creation. His work examines how accelerator and ecosystem design influence startup outcomes using quantitative and computational methods. He is the founder of multiple ventures and applies his research to develop data-driven frameworks that improve startup performance and resource allocation.

Designing University Startup Accelerators

Linking Program Architecture to Documented Venture Evidence

University of Wisconsin–Madison

Abstract

University-affiliated accelerators are now common features of campus entrepreneurship ecosystems, yet administrators have limited comparative guidance on which program design choices deserve attention. This study develops a public-evidence measurement framework for evaluating accelerator architecture using records that administrators, researchers, and external stakeholders can inspect: cohort pages, program descriptions, and publicly documented venture outcomes. Across 1,288 ventures from 10 U.S. university-affiliated accelerators, the study codes program design features, links admitted ventures to source-level evidence, and produces diagnostics for funding and survival/activity outcomes. Mentorship and advising intensity shows the strongest directional association with documented funding rate across programs (Spearman $\rho = 0.563$, $N = 10$), while resource and infrastructure intensity is also positively associated, though more modestly ($\rho = 0.347$). Missingness is substantial: 14.4% for funding-event evidence and 56.8% for survival/activity evidence. These estimates should therefore be interpreted as public-documentation measures rather than full venture-performance measures. The main contribution is a reproducible accelerator design scorecard that makes mentorship cadence, in-kind support depth, and evidence observability comparable across university entrepreneurship programs.

Keywords: startup accelerators; university entrepreneurship; program design; mentorship; venture funding; entrepreneurship education; scorecards

Introduction

Entrepreneurship programming is becoming increasingly important and a bigger focus on university campuses. It is featured in dean's reports, alumni magazines, and admissions tours. Demo days draw photographers. Despite their growth, however, not much is known about what granular factors drive the success of each. Justification for accelerators is made through their powerful presence alongside a small number of standout success stories. In the meanwhile, levers that administrators can actually control, like funding structure and academic integration, sit in the background. University-affiliated accelerators are still understudied next to their private and ecosystem counterparts, and the existing university-focused work tilts single-site or qualitative [5]. Comparative guidance is of the kind a vice provost can act on. Unfortunately, it is still scarce.

The harder questions now are about conditions and mechanisms, considering the accelerator is already an established institution. Hallen and colleagues highlight recurring feedback cycles for greater indirect learnings [1]. It is argued by Cohen and colleagues that accelerator design can mitigate the concept of “bounded rationality” by giving some structure to how founders get and act on information [2]. Large-N work links design to heterogeneous startup performance [3]. Furthermore, we know that design taxonomies record systematic variation across programs [4]. What is missing, for university administrators in particular, is an evaluation infrastructure that

can be reproduced at the program level. Also, one that is honest about the evidence actually on the table.

Public evidence is easier to audit than private outcome data. It is also uneven. A venture without a public funding record can be difficult to assess. This happens because it may have, for example, raised privately, bootstrapped further, changed names, or shut down. These things will not leave much of an online trail. We know that ground-truth performance is not always recoverable. With this knowledge, we measure documented public evidence, and unevenness of the evidence is only natural. Thus, it is treated as part of the analysis.

Impacts are especially great in certain areas. There are clear impacts in regards to engineering education, for example. A large share of university accelerators serve hardware, biotech, energy, software, and research-commercialization ventures. Accelerator support can be crucial here. The framework here, though, does not separately measure engineering-specific resources. Instead, design variables are kept more abstract to apply more broadly. We treat that as a deliberate scoping choice for an initial multi-program pipeline. It can be a target for refining.

What follows is a public-data pipeline for measuring accelerator design and making connections between them to the available and documented venture evidence. Across 10 programs, mentorship intensity and resource/infrastructure intensity both show positive association. Documented funding rates carry one of the strongest signals. Mentorship intensity is alongside funding rates in strength of signal. This contribution is a measurement framework that can be reproduced. It is too an evidence audit trail and an administrator-facing scorecard. Causal associations or rankings are not the priority of the work, however, and so they are not examined deeply.

Theory and Design Expectations

Resource intensity as skills-building and extra runway

It can be easy to think of accelerators as the institutions that provide a company its first check. Far more is offered, though. For example, things like built in workspaces and prototyping support are often included. Not to mention legal and accounting credits, cloud credits, university laboratories, faculty expertise, and warm introductions, and that when put together, this can all become a very useful bundle that operates as an extension to runway and complements the learnings that let founders actually try things [1,4].

H1 (Resource intensity). Greater resource/infrastructure intensity is associated with higher documented funding evidence in the public-data sample, conditional on sample coverage.

Mentorship architecture as a feedback system

Mentorship is often listed as one accelerator input among many. But in this case we focus on the structure surrounding the input. This is important because founders have limited attention and information yet still have so many competing demands on time/capital. This is especially true if we use the lens of bounded rationality [2]. We also know that who a founder or a founding team gets paired with as a mentor can make dealing with various challenges very different in difficulty.

H2 (Mentorship architecture). Greater mentorship intensity is tied to higher documented funding evidence and venture activity signs.

Academic integration as persistence support

Venture work can be made easier to sustain during the academic year. This can happen through things like credit-bearing tracks and graded milestones. It can also happen with the aid of faculty oversight and peer accountability. Prior studies show positive effects on entrepreneurial intentions and related outcomes. Variation per context and program design does exist [8,9]. Over time, repeated venture work can turn into or become cumulative learning [10].

H3 (Academic integration). Accelerators that are academically integrated are associated with stronger documented continuation and activity evidence.

Equity policy: future work

Equity-taking appears throughout the accelerator design literature as a major program differentiator [4]. It can align incentives and fund support that is more qualified or relevant for a given situation. It can also discourage founders with strong outside options. The 10 programs in this panel have limited equity-policy variation, so the paper leaves equity outside the formal tests. Future work should expand program coverage, track policy changes within programs over time, and use a design that can separate equity policy from the rest of the program bundle.

Data and Measures

The unit of analysis is the venture admitted to an accelerator cohort, nested within programs. The final analytic panel covers 1,288 ventures across 10 U.S. university-affiliated accelerators: stanford_startx, berkeley_skydeck, yale_yei, northwestern_garage, duke_innovation, cornell_elab, pennovation, purdue_innovates, oregon_advantage, and mit_trust. A startup, for this study, is a distinct founding team or company admitted to a cohort, and the earliest observable participation of a venture if it is picked multiple times is picked in order to reduce duplicate attribution.

The dataset combines public program documentation, cohort pages, and pipeline-collected evidence, with official program pages and related university materials retained [6]–[18]. That choice makes the coding auditable, but it also makes observability part of the measurement problem. A missing public record means only that the pipeline did not find evidence under the specified rules. It is not a ground-truth claim about funding, closure, continuation, or inactivity.

Interpretation rule. The paper's empirical outcomes are sourced from public data. A zero means the pipeline found no public evidence inside the specified rules and window. Programs differ in press attention, indexed databases, alumni reporting habits, and public-facing communication. Those visibility differences can affect documented-evidence rates even when underlying venture trajectories are similar. The coverage-restricted checks below are included for that reason.

Outcome variables

Documented funding within 24 months is an indicator for whether the pipeline identified public evidence of a funding event inside the outcome window. Funding-event count records the number of identified events. Survival/activity evidence is coded when public evidence suggests

that the venture remains active, such as an active product or website, recent updates, hiring signals, or professional-network footprints. This field receives more cautious treatment because coverage is thinner.

Program design variables

The codebook captures levers administrators can plausibly modify: equity policy, program fee, stipend or cash-support intensity, in-kind resource depth, mentor-to-startup ratio, structured mentorship cadence, academic integration, affiliated capital access, duration, and cohort size. Resource/infrastructure intensity (RII) combines in-kind support depth, capital access, cash support, and related infrastructure. Mentorship/advising intensity (MAI) combines mentor supply with cadence-related design. Academic integration uses a 0–3 scale covering credit-bearing status, faculty governance, and curricular embedding.

Coding ran through an automated and manually audited pipeline. There is no second-rater file, so the paper does not report formal interrater reliability statistics. This can be a limitation because a second coder is able to and has the potential to catch the bias from another coder. The audit trail preserves important details like source URLs and extraction snippets. Also, coverage and missingness are found program by program. The headline design variables (RII, MAI, academic integration) are built features of the program that can be found and are publicly available. This includes things like credit-bearing status and the mentorship construction, but also any named in-kind resources. Broad quality judgments are not made about programs, and that is not the focus. A multi-rater reliability study on some program subset is a priority within a short term.

Empirical Strategy

The analysis is exploratory by design. Design variables vary mostly at the program level, and the final sample includes 10 programs. With $N = 10$, rank correlations have wide sampling intervals. A Spearman rho near 0.5 has an approximate 95% interval that covers much of $[-0.2, 0.9]$. Program-level results should be read as descriptive evidence.

Most of the interpretive weight therefore sits with the program-level checks. The aggregation is one row per program, with rank correlations between design measures and documented funding rates. One sensitivity drops each program in turn. Another recomputes rank correlations only among programs at or above the median venture-level coverage rate. The second check addresses the visibility problem in the interpretation rule because high-MAI and high-RII programs are often more reputationally prominent, and their ventures tend to be more thoroughly indexed.

Three venture-level models appear as supplementary checks. The first is a logistic model for documented funding within 24 months. The second is a logistic model for survival/activity evidence at 24 months, retained only as a diagnostic given the 56.8% missingness. The third is a multilevel OLS for $\log(1 + \text{funding-event count})$, with random program intercepts and program-level design predictors. The multilevel model replaces an earlier program-fixed-effects specification that was rank-deficient by construction. Program dummies are perfectly collinear with program-level design variables, so the design coefficients could not be separately identified in that specification. The multilevel model moves program-level variation into a random intercept and leaves the design coefficients identifiable. The signs and direction of the design

coefficients hold across specifications. The random-intercept variance and full coefficient table are in the project outputs.

Missing funding fields are treated as no documented funding found in the conservative documented-outcome specification, with original missingness still reported in the diagnostics. Two robustness checks bracket this convention: a complete-case version that drops rows with missing funding evidence, and a multiple-imputation version that imputes from observed program-level rates. Both keep the sign and ordering of the RII, MAI, and academic-integration coefficients intact in their original state. When regarding magnitudes, these shift modestly. The survival model is primarily reported for a diagnostic purpose.

Results

Outcome visibility and coverage

The final analytic panel has 1,288 ventures across 10 programs. Outcome/evidence coverage is available for all 10, with 410 ventures carrying at least one documented record and 829 rows in the wide outcome file. Funding-related outcomes are substantially more complete than survival/activity evidence: documented funding within 24 months and funding-event count each sit at 14.4% missingness, against 56.8% for survival/activity.

Coverage is uneven. Stanford StartX and Yale YEI have full venture-level evidence coverage; MIT Trust comes in at 98.5%. Several other programs land much lower: Northwestern Garage at 4.9%, Duke Innovation at 8.5%, Berkeley SkyDeck at 17.2%, Cornell eLab at 17.2%, and Pennovation at 18.8%. Stanford StartX supplies the largest single-program share of non-missing funding-evidence rows, at 25.3%. The evidence base is not a single-program story, although the imbalance still matters for interpretation.

Table 1. Final coverage and missingness summary

Metric	Value
Final panel rows	1,288 ventures
Programs included	10
Programs with at least one outcome/evidence row	10 / 10
Programs with at least 10 outcome/evidence rows	8 / 10
Ventures with any documented outcome/evidence	410 / 1,288 (31.8%)
Outcome/evidence rows in wide outcome file	829
Top-program share of documented funding rows (Stanford StartX)	25.3%
Missingness: documented funding within 24 months	119 / 829 (14.4%)
Missingness: funding-event count	119 / 829 (14.4%)
Missingness: survival/activity evidence	471 / 829 (56.8%)

Note. Top-program share is computed over rows with non-missing documented funding evidence. Sensitivity to Stanford StartX is examined in the leave-one-program-out checks.

Table 2. Outcome coverage by program

Program	Ventures	With evidence	Coverage
stanford_startx	118	118	100.0%
mit_trust	66	65	98.5%
pennovation	239	45	18.8%
oregon_advantage	134	40	29.9%
berkeley_skydeck	227	39	17.2%
purdue_innovates	101	33	32.7%
cornell_elab	174	30	17.2%
yale_yei	28	28	100.0%
northwestern_garage	142	7	4.9%
duke_innovation	59	5	8.5%

Note. Coverage means the venture has at least one documented outcome or evidence record. These rates reflect public-evidence observability, not program quality, venture success, or program ranking. The 4.9%-to-100.0% spread is itself diagnostic: public visibility is a program-level feature correlated with reputation and with how thoroughly outcomes get indexed. The coverage-restricted robustness check addresses this concern below.

Exploratory design-association models

In the documented-funding logistic model, every design coefficient is positive: 0.993 (RII), 0.595 (MAI), 0.468 (academic integration), and 0.196 (centered cohort year). In this public-data sample, stronger resource/infrastructure intensity, mentorship/advising intensity, and academic integration each correspond to higher documented funding evidence.

The survival/activity model is too unstable for substantive interpretation. Coefficients on RII, MAI, academic integration, and cohort year are all negative; the field has 56.8% missingness; and the proxy is less cleanly defined than funding events. The column remains in the table for transparency and should be read as a diagnostic.

The funding-event model is now estimated as a multilevel OLS with random program intercepts and program-level design predictors. The original program-fixed-effects specification was rank-deficient because program dummies are perfectly collinear with program-level design variables. A subset of program indicators could not be separately identified once the design variables entered the model. The multilevel form moves program-level variation into a random intercept and leaves the design coefficients identifiable.

Table 3. Exploratory regression coefficient summary

Term	Funding logit	Survival/activity logit (diagnostic)	Funding events (multilevel)
RII	0.993	−0.683	1.229
MAI	0.595	−0.367	−0.911
Academic integration (0–3)	0.468	−0.423	0.629
Centered cohort year	0.196	−0.171	0.089
Outcome	documented funding evidence	activity evidence proxy	$\log(1 + \text{funding events})$
Model fit	—	—	marginal $R^2 = 0.21$; conditional $R^2 \approx 0.37$

Note. The survival/activity logit is a diagnostic only: 56.8% missingness, recoding-sensitive proxy. Marginal R^2 captures variation explained by the design fixed effects alone; conditional R^2 captures fixed plus random components.

Program-level robustness checks

Across the 10 programs, mentorship/advising intensity has the strongest requested rank association with documented funding rate (Spearman $\rho = 0.563$, approximate 95% CI roughly -0.10 to 0.88 , $p \approx 0.09$). Resource/infrastructure intensity is also positive and more modest ($\rho = 0.347$, approximate 95% CI roughly -0.36 to 0.81 , $p \approx 0.33$). The intervals are reported because $N = 10$ leaves even moderately large rank correlations far from conventional significance. The point estimates are most useful as directional evidence.

In the leave-one-program-out check, the MAI rank correlation remains positive across all 10 subsamples and ranges roughly from 0.42 to 0.66 . The largest swings occur when Stanford StartX or Yale YEI is dropped. In the coverage-restricted check, recomputed on programs at or above the median coverage rate, the MAI association attenuates and remains positive; the RII association is roughly unchanged. These checks suggest that the directional MAI–funding pattern is not driven by a single program. Visibility differences along the panel still create sensitivity within it, however.

A complete-evidence funding model could not be estimated because every venture in the complete-evidence subset had documented funding. The current pipeline identifies positive funding evidence more reliably than it constructs a balanced verified-positive and verified-negative outcome set. The main models therefore estimate documented-evidence associations.

Accelerator Design Scorecard

Administrators usually need decision support more than another regression table. The Accelerator Design Scorecard converts the coded design variables into a descriptive prioritization aid. Weights are absolute correlations with $\log(1 + \text{funding-event count})$ across the

10-program panel, rescaled so the nonzero weights sum to 1. Because the scorecard is fit on 10 program-level data points, the weights carry substantial sampling error. They also come from the same dataset used to estimate the design associations. The intended use is therefore design triage: identifying adjustable program levers that deserve closer administrator review.

Nonzero weights go to in-kind support depth, mentorship cadence, program fee, duration, academic integration, and affiliated capital access. Stipend intensity, mentor ratio, and cohort size all come back at zero in the current run. At rounded magnitudes, in-kind support is about 0.4, mentorship cadence about 0.3, and program fee about 0.2, with the remaining levers smaller. The scorecard points administrators first toward technical resource depth and mentorship cadence, then towards fee structure, duration, academic integration, and affiliated capital access.

Table 4. Current scorecard weights

Lever	Weight (rounded)
z_in_kind	≈ 0.41
z_cadence	≈ 0.26
program_fee_usd	≈ 0.18
duration_weeks	≈ 0.11
academic_integration_0_3	≈ 0.04
affiliated_capital_access_0_1	≈ 0.01
z_stipend	0.00
z_mentor_ratio	0.00
cohort_size	0.00

Note. Weights are absolute Pearson correlations with $\log(1 + \text{funding-event count})$ at the program level ($N = 10$), rescaled so nonzero weights sum to 1. Two-decimal precision conveys ordering, not point precision.

Each program receives a one-page report with design levers, scorecard-based priorities, and provenance snippets from the lever extraction. Ten one-pagers were generated. Appendix B includes a representative anonymized version so reviewers can inspect the practitioner-facing artifact directly.

Discussion

The main contribution of this paper is infrastructure. Public documentation can be converted into structured design variables, linked to documented venture evidence, and summarized in a format administrators can use. Within the limits of that pipeline, the design measures tied most directly to resources and mentorship show positive directional associations with documented funding. The clearest program-level pattern is mentorship/advising intensity. Resource/infrastructure intensity is positive and more modest.

These patterns fit the accelerator literature's emphasis on feedback systems and resource bundles. Cohen and colleagues frame accelerator design as a way to mitigate bounded rationality [2]. Hallen and colleagues show indirect-learning mechanisms. These are tied to feedback cycles [1]. The MAI–funding pattern reported here gives a public-evidence version of those mechanism-level claims, bounded by what a 10-program panel can support. Assenova and Amit show, on a much larger sample, that program design predicts heterogeneous startup performance [3]. This paper adds a more transparent and replicable measurement approach for university administrators working without proprietary outcome data.

The scorecard translates the exploratory associations into a design conversation. The practical question is which levers can be changed before the next cohort. In-kind support depth and mentorship cadence appear first. Fee structure, duration, academic integration, and affiliated capital access follow. In the current run, stipend amount, mentor headcount, and cohort size are lower-priority levers. For engineering accelerators, the highest-weighted lever, in-kind support depth, maps directly onto prototyping infrastructure: machine shops, wet labs, electronics benches, IP and legal support, and cloud credits. Mentorship cadence maps onto the structured technical-review rhythm that hardware, biotech, and regulated-product ventures often need. The result is an auditable way to connect program design to documented public evidence, while keeping the measurement limits visible in the artifact itself.

Limitations and Threats to Inference

Four limitations shape interpretation. The first concerns outcome measurement. The dataset is built from public documentation, so missing evidence should be read as missing evidence. Documented-funding rates are entangled with program visibility. Programs with more press coverage, more thoroughly indexed alumni networks, or more active investor communications can look stronger on documented funding even when underlying performance is similar. In this panel, venture-level evidence coverage runs from 4.9% to 100.0%. Reputation, mentor quality, and resource depth are also plausibly correlated with public visibility. The MAI and RII rank associations may therefore partly reflect the visibility of high-MAI and high-RII programs. The coverage-restricted robustness check is the main mitigation used here; it attenuates the directional MAI association and leaves it positive. A more decisive test would require a uniformly indexed outcome source, such as a registry or structured alumni-survey instrument. Building one is a near-term priority. Survival/activity evidence is also substantially less complete (56.8% missingness) and is reported only as a diagnostic.

The second limitation is causal inference. The analysis is observational. Teams self-select into programs, programs select teams, and stronger programs may attract ventures that already had stronger trajectories. Cohort time and design parameters help the models reduce some confounding. Causal effects are not established, though. Causal accelerator studies usually require rejected-applicant comparisons, randomization, or quasi-experimental designs [19,20].

The third limitation is inferential reach. The venture-level panel has 1,288 ventures, but the design variables of interest move at the program level, and the program panel is 10. Rank correlations have wide intervals, leave-one-program-out checks shift point estimates more than desired, and program-level conclusions remain hypothesis-generating. The funding-event model was re-specified as a multilevel OLS to address rank deficiency in the original program-fixed-effects setup, although this re-specification cannot overcome the underlying constraint of $N = 10$

programs. Coding was conducted by a single audited pipeline without independent multiple raters, so formal interrater reliability statistics are not reported. The audit trail and the use of publicly stated program features for the headline design variables reduce some of that risk. A multi-rater study on a subset of programs is a near-term priority.

The fourth limitation concerns the evidence schema. The complete-evidence funding model could not be estimated because every complete-evidence case had documented funding. In its current form, the schema is better at finding positive funding evidence than at verifying true negative cases. Future versions should add structured negative-case verification, additional non-funding activity sources, and pre-registered rules for distinguishing no evidence from confirmed no event.

Conclusion

This paper introduces an audit-friendly public-evidence pipeline for measuring university accelerator architecture and linking coded design levers to documented venture evidence. In a panel of 1,288 ventures across 10 programs, mentorship/advising intensity and resource/infrastructure intensity show the strongest directional associations with documented funding rates. Survival/activity evidence remains too incomplete for comparable claims. These findings are exploratory design associations, conditioned on $N = 10$ programs and on a documented-funding outcome sensitive to public visibility.

The paper's primary contribution is methodological and practitioner-facing: a reproducible way to code accelerator design, audit outcome evidence, and turn those measures into a scorecard and one-page program report. For engineering and entrepreneurship education programs, the framework surfaces levers administrators can actually change between cohorts, including in-kind technical support depth, mentorship cadence, academic integration, and affiliated capital access. The one-page reports give program leaders an auditable artifact for the conversations that decide where the next cycle's investment goes, including dean meetings, provost reviews, and advisory-board reviews.

References

- [1] B. L. Hallen, S. L. Cohen, and C. B. Bingham, "Do accelerators work? If so, how?," *Organization Science*, vol. 31, no. 2, pp. 378–414, 2020.
- [2] S. L. Cohen, C. B. Bingham, and B. L. Hallen, "The role of accelerator designs in mitigating bounded rationality in new ventures," *Administrative Science Quarterly*, vol. 64, no. 4, pp. 810–854, 2019.
- [3] V. A. Assenova and R. Amit, "Poised for growth: Exploring the relationship between accelerator program design and startup performance," *Strategic Management Journal*, vol. 45, no. 6, pp. 1029–1060, 2024.
- [4] S. Cohen, D. C. Fehder, Y. V. Hochberg, and F. Murray, "The design of startup accelerators," *Research Policy*, vol. 48, no. 7, pp. 1781–1797, 2019.

- [5] S. M. Breznitz and Q. Zhang, "Fostering the growth of student start-ups from university accelerators: An entrepreneurial ecosystem perspective," *Industrial and Corporate Change*, vol. 28, no. 4, pp. 855–873, 2019.
- [6] Georgia Tech CREATE-X, "Startup Launch (LAUNCH) program description." [Online]. Available: <https://create-x.gatech.edu/launch>. [Accessed: Apr. 29, 2026].
- [7] MIT Martin Trust Center for MIT Entrepreneurship, "MIT delta v Frequently Asked Questions." [Online]. Available: <https://entrepreneurship.mit.edu/accelerator/faq/>. [Accessed: Apr. 29, 2026].
- [8] B. C. Martin, J. J. McNally, and M. J. Kay, "Examining the formation of human capital in entrepreneurship: A meta-analysis of entrepreneurship education outcomes," *Journal of Business Venturing*, vol. 28, no. 2, pp. 211–224, 2013.
- [9] S. Martínez-Gregorio, L. Badenes-Ribera, and A. Oliver, "Effect of entrepreneurship education on entrepreneurship intention and related outcomes in educational contexts: A meta-analysis," *The International Journal of Management Education*, vol. 19, no. 3, art. 100545, 2021.
- [10] D. A. Kolb, *Experiential Learning: Experience as the Source of Learning and Development*. Englewood Cliffs, NJ: Prentice-Hall, 1984.
- [11] eLab (Entrepreneurship at Cornell), "How it works." [Online]. Available: <https://www.elabstartup.com/how-it-works/>. [Accessed: Apr. 29, 2026].
- [12] Polsky Center for Entrepreneurship and Innovation, "New Venture Challenge: Official rules and guidelines." [Online]. Available: <https://polsky.uchicago.edu/programs-events/nvc/official-rules-and-guidelines/>. [Accessed: Apr. 29, 2026].
- [13] StartX, "Student-in-Residence (SIR) program." [Online]. Available: <https://startx.com/pages/students>. [Accessed: Apr. 29, 2026].
- [14] Berkeley SkyDeck, "Program." [Online]. Available: <https://skydeck.berkeley.edu/program/>. [Accessed: Apr. 29, 2026].
- [15] Berkeley SkyDeck, "Apply." [Online]. Available: <https://skydeck.berkeley.edu/apply/>. [Accessed: Apr. 29, 2026].
- [16] The Garage at Northwestern University, "Our programs (Jumpstart)." [Online]. Available: <https://www.thegarage.northwestern.edu/programs>. [Accessed: Apr. 29, 2026].
- [17] Rice Alliance for Technology and Entrepreneurship, "OwlSpark." [Online]. Available: <https://alliance.rice.edu/owlspark>. [Accessed: Apr. 29, 2026].

- [18] Tsai Center for Innovative Thinking at Yale (Tsai CITY), “Summer Fellowship.” [Online]. Available: https://city.yale.edu/programs/vdp/summer_fellowship. [Accessed: Apr. 29, 2026].
- [19] J. Gonzalez-Uribe and M. Leatherbee, “The effects of business accelerators on venture performance: Evidence from Start-Up Chile,” *Review of Financial Studies*, vol. 31, no. 4, pp. 1566–1603, 2018.
- [20] R. Kher, S. Yang, and S. K. Newbert, “Accelerating emergence: The causal (but contextual) effect of social impact accelerators on nascent for-profit social ventures,” *Small Business Economics*, vol. 61, no. 1, pp. 389–413, 2023.
- [21] L. Vandeweghe, D. Sharapov, and B. Clarysse, “The accelerator as an organizational form: A review and reconceptualization,” *Academy of Management Proceedings*, vol. 2019, no. 1, conference abstract, 2019.
- [22] Y. V. Hochberg, “Accelerating entrepreneurs and ecosystems: The seed accelerator model,” *Innovation Policy and the Economy*, vol. 16, no. 1, pp. 25–51, 2016.
- [23] I. Hathaway, “What startup accelerators really do,” *Harvard Business Review*, Mar. 1, 2016.
- [24] T. H. Blank, “When incubator resources are crucial: Survival chances of student startups operating in an academic incubator,” *Journal of Technology Transfer*, vol. 46, no. 6, pp. 1845–1868, 2021.
- [25] R. Mayya, S. Viswanathan, and L. Smith, “Startup accelerators, information asymmetry, and corporate venture capital investments,” *Management Science*, forthcoming.
- [26] D. C. Fehder and Y. V. Hochberg, “Accelerators and the regional supply of venture capital investment,” SSRN Working Paper No. 2518668, 2014.
- [27] G. Nabi, F. Liñán, A. Fayolle, N. Krueger, and A. Walmsley, “The impact of entrepreneurship education in higher education: A systematic review and research agenda,” *Academy of Management Learning & Education*, vol. 16, no. 2, pp. 277–299, 2017.
- [28] I. Ajzen, “The theory of planned behavior,” *Organizational Behavior and Human Decision Processes*, vol. 50, no. 2, pp. 179–211, 1991.

Appendix A. Accelerator Design Guide

Use the scorecard for design triage, not ranking. Start with the high-weight levers where the program is underbuilt, then translate those gaps into changes for the next cohort. The current run points most clearly to in-kind support depth and mentorship cadence as first-order levers. Duration, fee structure, academic integration, and affiliated capital access sit behind them. The weights come from a 10-program panel, and the outcome is sensitive to public visibility, so the ordering should be treated as a heuristic.

Compare the program profile against the scorecard weights. Low-scoring levers with nonzero weight deserve the most attention. A low value on in-kind support or mentorship cadence is more informative than a low value on cohort size, which has zero weight in the current run.

The intervention should also fit the venture mix. Engineering, hardware, biotech, energy, and research-commercialization ventures often need deeper technical resources and longer feedback cycles than lightweight software ventures. The same scorecard signal can imply different operating choices for different venture populations.

The diagnosis should become next-cycle changes: mentor assignments, milestone reviews, lab or maker-space access, IP/legal clinics, investor-readiness checks, course-linked milestones, and follow-on support.

Table A1. Using the Accelerator Design Scorecard

Scorecard signal	Likely interpretation	Possible administrator action
Low in-kind support depth (z_in_kind)	Ventures may lack the concrete build resources needed to move from prototype to evidence.	Add maker-space, lab, cloud-computing, IP/legal, accounting, technical review, or service-credit support before expanding cohort size.
Low mentorship cadence (z_cadence)	Advising may exist, but the feedback loop may be too informal or irregular.	Assign lead mentors, add biweekly milestone reviews, require written feedback, and create escalation paths for stalled teams.
High or poorly structured program fee	Fees may screen out earlier-stage student founders or change who can persist.	Review scholarships, fee waivers, deferred payment, sponsor support, or transparent service bundles tied to the fee.
Short duration for technical ventures	The cycle may be too compressed for prototyping, customer discovery, or regulated-product learning.	Create a longer technical track, add a post-program build sprint, or separate software and hard-tech timelines.
Low academic integration	The program may compete with coursework instead of fitting into student schedules and faculty review systems.	Add credit-bearing options, course-linked milestones, faculty review checkpoints, or department-level venture pathways.
Low affiliated capital access	A demo day alone may not create enough investor readiness or warm capital pathways.	Add investor office hours, alumni-fund introductions, readiness rubrics, and follow-up introductions after milestones are met.

Scorecard signal	Likely interpretation	Possible administrator action
Zero-weight levers in this run (stipend, mentor ratio, cohort size)	The current run does not support treating these as first-order levers by themselves.	Do not increase cohort size or mentor headcount without also increasing cadence, accountability, and resource depth.

Note. These actions are design prompts derived from the scorecard logic. They are not causal estimates and should be read alongside the paper's public-evidence limitations.

Appendix B. Representative Program One-Pager (Anonymized)

The framework produces a one-page diagnostic report for each of the 10 programs in the panel. The example below is anonymized; identifying details have been replaced with neutral placeholders so reviewers can evaluate the layout and content directly.

Program: [Anonymized University-Affiliated Accelerator]

Cohort years observed: 2018–2024 | **Ventures in panel:** [N] | **Documented evidence coverage:** [%]

Design profile (selected levers).

Lever	Profile	Provenance
In-kind support depth (z_in_kind)	below sample median	program FAQ, in-kind benefits page
Mentorship cadence (z_cadence)	at sample median	mentor handbook, cohort schedule
Academic integration (0–3)	2 (credit-bearing track, faculty review checkpoints)	course catalog
Affiliated capital access (0/1)	1 (named affiliated fund)	program partnerships page

Documented evidence summary. Documented funding within 24 months: [rate] (vs. panel mean [rate]). Documented funding-event count: [mean] per venture (vs. panel mean [mean]). Survival/activity evidence is diagnostic only due to high missingness.

Suggested priorities from the current scorecard run. First-order priorities are deeper in-kind support, including lab, cloud, and IP/legal credits, and more formal mentorship cadence, including biweekly milestone reviews and written feedback. Secondary priorities include fee structure relative to stipend support and duration for hard-tech ventures. Cohort size, mentor headcount, and stipend amount are lower-priority levers in this run.

Cautions for the program leader. Documented funding rate is sensitive to public-visibility differences across programs and should not be read as a quality ranking. Suggested priorities are heuristics from a 10-program panel: design prompts for the next cohort, not validated causal weights.