

國立臺灣大學管理學院資訊管理學研究所



碩士論文

Department of Information Management

College of Management

National Taiwan University

Master's Thesis

新創早期成功預測中偏見消除方法的比較研究

From Bias to Balance: A Comparative Study of Bias  
Mitigation Methods in Startup Early Success Prediction

翁子婷

Tzu-Ting Weng

指導教授：魏志平 博士

Advisor: Chih-Ping Wei, Ph.D.

中華民國 114 年 7 月

July 2025

## 致謝



兩年碩士時光如此快速地結束了，在這段期間，我深刻體會到知識就像是一種悖論，當我知道的越多，就越能發現有更多未知的領域等著我去探索。而在茫茫的學海中，魏老師就像一座燈塔，指引著我們前進的方向，能成為魏老師的學生真的是非常幸運的一件事。這一路上，從一開始試圖看懂一篇論文，到最後能夠自己寫出一篇論文，都要感謝魏老師的細心指導，即便我們實驗室成員的研究方向各自不同，魏老師總能在各種領域提出新的想法並給予具體建議，幫助我們突破視野，從更多角度去分析、思考與驗證自己的研究。

還記得當初和老師確定研究題目時，因為過去的學長姐幾乎沒有人做演算法公平問題的研究，讓我一度感到猶豫。但老師說，這個題目除了我自己能從中學學習外，老師本身也能因此學到新的東西，讓我感受到老師真的是一位非常好學、願意擁抱新知的人。除了學術之外，老師在為人處事上也讓我學到很多，無論是對學生的尊重與包容、願意花時間傾聽我們的想法，還是遇事時的冷靜思考與穩重判斷，都讓我深受影響。

最後，我也要感謝實驗室夥伴們，能和大家一起互相激勵、成長，最終一同完成學業，真的是一件很令人感動又珍貴的事。感謝思涵學姊陪我和老師一起開會，提出各種想法，也在論文寫作上提供許多指導；感謝沛原學長在學期間關心我的論文進度，並給我許多關於未來職涯的寶貴建議；感謝虹鈞姐為我們打點實驗室的大小事，讓我們能順利準備口試。最後，我要感謝我的家人，在我陷入低潮時給予關心，幫我處理生活上的瑣事，讓我能無後顧之憂地繼續前行。

一路上我受到許多人的幫助，希望未來的我，也能成為推動他人的一股助力。

翁子婷 謹識

于台灣大學資訊管理學研究所

中華民國一百一十四年八月

## 摘要



隨著創業投資家 (VC) 越來越依賴於使用機器學習來輔助他們的投資決策，這些演算法是否會延續過去募資結果中所存在的歧視性偏見，成為了一項令人關心的議題。這些偏見大多起因於新創公司早期缺乏足夠的可量化資料，使得投資人進行投資決策時，往往會依賴他們對創辦人團隊的主觀判斷，而這很有可能招致有關人口統計上的刻板印象與歧視。為了避免這類偏見在演算法中被進一步強化，確保決策系統中的公平性是避免創業環境下的資源錯誤分配、以及平等資金機會的關鍵。

在本研究中，針對新創早期成功預測任務，我們考量了三種常見的潛在歧視來源，包含地理區域、創辦人性別以及種族，並且實作與比較了三種公平性方法：特徵遮蔽 (feature-blind)、正則化法 (regularization-based) 與梯度反轉 (gradient reversal)。這些方法皆可處理具有混合資料型態的多個敏感屬性 (sensitive attributes)。我們的實驗結果顯示，儘管提升公平性會略微影響到目標任務的預測效能，但正則化法與梯度反轉法皆能有效改善模型公平性。

除了比較模型表現外，本研究也進一步識別出哪些子群體最容易受到模型偏見影響，例如創辦人中女性比例超過 75% 的新創企業是最不受基準模型的青睞的。我們也分析了哪個敏感屬性造成了最多的模型偏見。這些研究成果可為新創公司與創投提供實務上的見解，對新創企業而言，採用具公平措施的模型能提升他們平等地獲得資金的機會，而不受既有歧視的影響，進而打造更具包容性的創業環境；對投資人而言，這些模型有助於幫助他們發掘那些可能因偏見而被忽略的投資機會，並建立更平衡的投資組合。

**關鍵字：**公平性機器學習、新創公司成功預測、新創公司分析、演算法偏見、演算法公平、預測建模、表徵學習、決策支援系統

# Abstract



As venture capital (VC) firms increasingly adopt machine learning (ML) tools to support investment decisions, concerns arise regarding the potential perpetuation of historical biases embedded in past funding outcomes. These biases often stem from the limited availability of quantifiable data on early-stage startups. As a result, investment decisions depend heavily on subjective assessments of founding teams, which introduces risks of demographic stereotyping and discrimination. To prevent the reinforcement of such biases, ensuring fairness in ML-based decision systems is therefore critical to mitigating systematic resource misallocation and promoting equitable access to capital.

This study investigates fairness-aware startup early success prediction by examining three commonly cited sources of potential biases, including geographic region, founder gender, and race. We implement and compare three fairness-aware approaches: feature-blind, regularization-based, and gradient reversal, each capable of handling multiple sensitive attributes of mixed data types. Our empirical results demonstrate that, while introducing modest trade-off in predictive performance, both the regularization and gradient reversal methods effectively enhance fairness.

Beyond performance evaluation, this study identifies subgroups most impacted by model biases, such as startups with over 75% female founders, and highlights which

sensitive attribute contributes most to observed disparities. The findings offer actionable insights for both startups and VC practitioners. For startups, the adoption of fairness-aware methods can improve fairer access to funding opportunities and foster a more inclusive entrepreneurial landscape. For investors, these methods may help uncover overlooked ventures and support more balanced portfolio construction.

**Keywords:** Fairness-aware machine learning, Startup success prediction, Startup analytics, Algorithmic bias, Algorithmic fairness, Predictive modeling, Representation learning, Decision support systems

# Table of Contents



|                                                                         |      |
|-------------------------------------------------------------------------|------|
| 致謝.....                                                                 | i    |
| 摘要.....                                                                 | ii   |
| Abstract.....                                                           | iii  |
| Table of Contents .....                                                 | v    |
| List of Figures.....                                                    | vii  |
| List of Tables .....                                                    | viii |
| Chapter 1 Introduction.....                                             | 1    |
| 1.1 Background .....                                                    | 1    |
| 1.2 Research Motivation.....                                            | 4    |
| 1.3 Research Objectives .....                                           | 5    |
| Chapter 2 Literature Review .....                                       | 7    |
| 2.1 Predictive Features for Startup Early Success Prediction.....       | 7    |
| 2.2 Existing Studies on Mitigating Unfairness in Machine Learning ..... | 10   |
| 2.2.1 Pre-Processing Methods .....                                      | 10   |
| 2.2.2 In-Processing Methods .....                                       | 11   |
| 2.2.3 Post-Processing Methods .....                                     | 14   |
| 2.3 Summary and Limitations of Existing Literature .....                | 15   |
| Chapter 3 Methodology .....                                             | 20   |
| 3.1 Definition of Startup Early Success .....                           | 20   |
| 3.2 Predictive Features Used in Our Research .....                      | 21   |

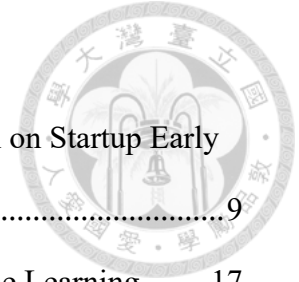
|                                                                                     |           |
|-------------------------------------------------------------------------------------|-----------|
| 3.3 Fairness-Aware Modeling Approaches .....                                        | 25        |
| 3.3.1 <i>Blind Method</i> .....                                                     | 25        |
| 3.3.2 <i>Regularization Method</i> .....                                            | 26        |
| 3.3.3 <i>Gradient Reversal Method</i> .....                                         | 27        |
| <b>Chapter 4 Experiments.....</b>                                                   | <b>33</b> |
| 4.1 Data Collection.....                                                            | 33        |
| 4.2 Baseline and Model Settings .....                                               | 34        |
| 4.3 Evaluation Design .....                                                         | 36        |
| 4.3.1 <i>Fairness Evaluation Metrics</i> .....                                      | 36        |
| 4.3.2 <i>Evaluation Procedure and Performance Metrics</i> .....                     | 40        |
| 4.4 Evaluation Results.....                                                         | 42        |
| 4.4.1 <i>Fairness-aware Methods</i> .....                                           | 42        |
| 4.4.2 <i>Group-wise Disparities across Sensitive Attributes</i> .....               | 45        |
| 4.4.3 <i>Attribution of Prediction Unfairness to Sensitive Attributes</i> .....     | 49        |
| 4.4.4 <i>Impact of Fairness Interventions on A Single Sensitive Attribute</i> ..... | 50        |
| <b>Chapter 5 Conclusion .....</b>                                                   | <b>53</b> |
| 5.1 Conclusion.....                                                                 | 53        |
| 5.2 Future Research Directions .....                                                | 55        |
| <b>References.....</b>                                                              | <b>57</b> |
| <b>Appendix.....</b>                                                                | <b>64</b> |

## List of Figures

|                                                                                   |    |
|-----------------------------------------------------------------------------------|----|
| Figure 1: Architecture of Gradient Reversal Approach for Fair Classification..... | 14 |
| Figure 2: Architecture of Our Gradient Reversal Approach for Fair Startup Success |    |
| Prediction.....                                                                   | 28 |



## List of Tables



|                                                                                                           |    |
|-----------------------------------------------------------------------------------------------------------|----|
| Table 1: Overview of Predictive Features Used in Prior Research on Startup Early Success Prediction ..... | 9  |
| Table 2: Summary of Existing Approaches to Fairness in Machine Learning .....                             | 17 |
| Table 3: Comparison of Fairness Methods in Handling Multiple and Mixed Type Sensitive Attributes.....     | 18 |
| Table 4: Model Settings of Baseline and Fairness-Aware Approaches .....                                   | 34 |
| Table 5: FairGap Computation Example Using Gender as Sensitive Attribute .....                            | 38 |
| Table 6: Evaluation Results of Baseline and Fairness-Aware Methods .....                                  | 41 |
| Table 7: Positive Prediction Rates by Region Across Methods .....                                         | 46 |
| Table 8: Positive Prediction Rates by Gender Composition Across Methods.....                              | 47 |
| Table 9: Positive Prediction Rates by Racial Composition Across Methods .....                             | 48 |
| Table 10: Impact of Isolated Sensitive Attributes on Fairness Metrics .....                               | 49 |
| Table 11: Impact of Fairness Interventions on Region Only .....                                           | 50 |
| Table 12: Impact of Fairness Interventions on Gender Only .....                                           | 51 |
| Table 13: Impact of Fairness Interventions on Race Only .....                                             | 51 |
| Table A1: Overview of Predictive Features Used in Our Methods.....                                        | 64 |
| Table A2: Descriptive Statistics of Our Dataset by Startup Success Label.....                             | 70 |

# Chapter 1 Introduction



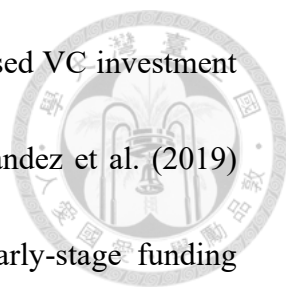
## 1.1 Background

Securing funding across various stages is a critical milestone in the lifecycle of startups. In particular, early-stage venture funding, such as Series A, is regarded as a key inflection point, as it often marks the transition from early development to scalable growth (Hellmann & Puri, 2000). In this stage, venture capital (VC) firms play a crucial role by providing not only financial resources but also strategic guidance to help startups foster sustainable expansion (Bygrave & Timmons, 1992). However, the path to success remains highly uncertain: only about one in three startups achieves initial profits within six years of founding (Reynolds, 2016). To mitigate the inherent risk associated with early-stage startup investments, venture capital firms are increasingly relying on machine learning (ML) to support their investment decisions (Astebro, 2021). Despite this trend, concerns have been raised that these data-driven systems may inherit and even amplify historical biases embedded in past VC decisions, potentially reproducing patterns of discrimination (Mehrabi et al., 2021).

The complexity and ambiguity of assessing nascent companies often give rise to biases in VC decision-making. Unlike established firms, the evaluation of early-stage startups often involves making high-stakes decisions under extreme uncertainty, as quantifiable information, such as financial records or product traction, is typically

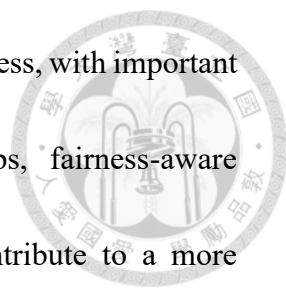
unavailable (Wei et al., 2025). Consequently, investment decisions tend to rely heavily on subjective assessments of founding teams (Corea et al., 2021). This subjectivity is further compounded by the presence of complex social signals and information asymmetry, which make objective evaluation particularly challenging in such ambiguous contexts (Gompers & Lerner, 2004). To overcome these obstacles, venture capitalists often resort to heuristics (Dale, 2015) and stereotypes (Bodenhausen & Wyer, 1985) to simplify judgment and improve decision-making efficiency, which may introduce or reinforce potential unfairness in funding outcomes.

Various manifestations of bias have been documented in VC investment decisions. Gender bias is one notable example, in which male investors often demonstrate less interest in female entrepreneurs compared to equally qualified male counterparts (Ewens & Townsend, 2020). Racial disparities have also been observed, with Black-owned startups receiving significantly less external equity funding at founding than White-owned firms (Fairlie et al., 2022; Paglia & Harjoto, 2014). Age bias further skews funding decisions, as investor evaluations follow an inverted-U pattern, favoring founders in a perceived optimal middle-age range (Matthews et al., 2024). Additionally, geographic bias plays a role, where startups located within a VC's home country are more likely to receive funding, while those based abroad are often overlooked (Coval & Moskowitz, 1999).



In addition to the unfair treatment faced by entrepreneurs, biased VC investment decisions also carry consequences for investors themselves. Hernández et al. (2019) conduct interviews with seven venture capitalists involved in early-stage funding decisions. Most of these investors actively seek to address biases, such as gender stereotypes, by expanding their own networks and diversifying the pool of promising entrepreneurs beyond familiar circles. These efforts are not solely driven by ethical considerations; rather, they also reflect a growing awareness of the potential financial impact of unfair practices. For example, empirical evidence shows that women-led firms can outperform male-led counterparts under comparable conditions (Gazanchyan et al., 2017), suggesting that gender bias may lead investors to overlook high-potential opportunities, ultimately resulting in suboptimal investment outcomes. Nonetheless, current fairness initiatives in venture capital predominantly focus on human decision-making, which may lack consistency and efficiency in mitigating biases (Hernandez et al., 2019). Moreover, limited attention has been paid to incorporating fairness considerations into machine learning systems that support VC decision-making (Te et al., 2023a).

When machine learning models are trained on biased historical data, they risk perpetuating or even exacerbating existing discrimination against underrepresented groups (Mehrabi et al., 2021). Incorporating fairness into ML systems is therefore

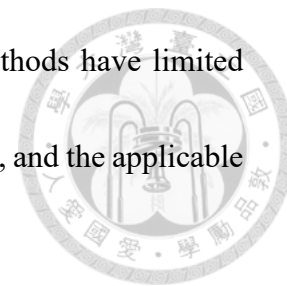


critical to mitigating such risks in the evaluation of startup early success, with important significance for both entrepreneurs and investors. For startups, fairness-aware approaches can enable more equitable access to capital and contribute to a more inclusive entrepreneurial ecosystem (Balachandra, 2020; Ivanitzki & Rashida, 2023). For venture capitalists, such research may help them uncover high-potential ventures that might otherwise be overlooked due to systemic biases, thereby enhancing both the diversity and potential returns of their investment portfolios (Hernandez et al., 2019). Furthermore, by integrating fairness considerations into ML models, investors are better positioned to make more informed and impartial decisions, reducing the risk of unlawful discrimination and promoting compliance with both legal requirements and ethical standards (Kumar et al., 2022).

## **1.2 Research Motivation**

Despite the increasing interest in algorithmic fairness, most existing methods are designed for general-purpose benchmark datasets that focus on individual-level classification tasks, such as credit scoring or income level prediction, and typically consider only one binary sensitive attribute at a time. However, these settings fail to capture the complexities of startup early success prediction, which requires evaluating fairness at the team level and handling multiple sensitive attributes that are often continuous or multi-label in nature. Such characteristics make the direct application of

existing fairness methods challenging in this domain, as most methods have limited capability to handle multiple sensitive attributes of mixed data types, and the applicable ones often require additional adaptation.

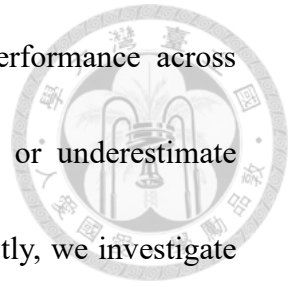


Moreover, prior fairness studies rarely explore performance metrics beyond accuracy, even though metrics such as precision and recall hold particularly meaningful implications in the context of early-stage funding prediction. As Moriarty et al. (2019) note, from the investor's perspective, incorrectly recommending an unpromising startup may lead to direct financial loss, while failing to recommend a promising one merely results in a missed opportunity. This rationale underscores why many investor-oriented systems prioritize optimizing precision. However, from the startup's standpoint, a lower recall of the prediction model can be far more damaging, as being wrongly rejected leads to lost funding opportunities and stunted growth. These dual perspectives suggest that fairness evaluations should not only focus on overall accuracy but also consider how different metrics disproportionately impact stakeholders.

### **1.3 Research Objectives**

In this study, we aim to explore algorithmic fairness in the context of startup early success prediction. Specifically, we implement and compare several fairness-enhancing methods, namely feature-blind, regularization-based training, and gradient reversal, to assess their effectiveness in mitigating biases in startup early success prediction. We

further evaluate the trade-off between fairness and predictive performance across different methods, and examine their tendencies to overestimate or underestimate certain types of startups by analyzing the prediction outcomes. Lastly, we investigate which sensitive attribute most significantly contributes to unfairness in prediction outcomes, thereby providing insights that may bring more equitable and responsible investment practices for both startup founders and venture capitalists.



## Chapter 2 Literature Review



### 2.1 Predictive Features for Startup Early Success Prediction

Existing literature commonly classifies features used for startup early success prediction into three major categories: company, founders, and investment. Each reflects different dimensions of a startup's potential in securing early-stage funding.

Company-level features describe the foundational and externally observable attributes of the startup. These include basic characteristics such as the company's geographic location (e.g., country and city), industry sector, age since founding, the presence of experienced advisors, number of launched products, and textual descriptions of the business (Krishna et al., 2016; Sharchilev et al., 2018; Te et al., 2023b). Although not widely adopted, one study has explored the use of media and public attention indicators, such as the number of news mentions, breadth of domain coverage, and topic modeling derived from Latent Dirichlet Allocation (LDA) applied to news content (Sharchilev et al., 2018).

Founder-level features focus on the human capital embedded in the entrepreneurial team. Commonly used variables include the number of founders, demographic attributes, educational background (e.g., degree level and institution ranking), and prior work experience, which serve as proxies for individual capabilities and team diversity and are widely recognized for their impact on startup performance (Sharchilev et al.,

2018; Te et al., 2023b).



Investment-related features reflect the startup's financing dynamics and the quality of external support. These can be further divided into two subcategories. The first pertains to funding details, such as total capital raised, number of past funding rounds, average time interval between rounds, burn rate (i.e., spending speed), and capital concentration (i.e., dependency on key investors). These features have been explored across multiple studies (Krishna et al., 2016; Sharchilev et al., 2018; Te et al., 2023b; Wei et al., 2025). The second subcategory encompasses investor characteristics, including the number and types of past investors (i.e., individuals or organizations). Some studies also consider investors' average centrality within the historical co-investment network and their prior success rates, capturing both the breadth and strength of investor backing (Wei et al., 2025).

Table 1 summarizes representative features from prior studies. These categories form the basis of the feature design in our proposed methods, which integrates both static and dynamic signals to enable a more comprehensive prediction of early-stage startup funding success.

**Table 1: Overview of Predictive Features Used in Prior Research on Startup**

| Early Success Prediction |                 |                                                                                 |                          |
|--------------------------|-----------------|---------------------------------------------------------------------------------|--------------------------|
| Category                 | Subcategory     | Features                                                                        | Source                   |
| Company                  | Basic features  | Location                                                                        | Krishna et al. (2016)    |
|                          |                 | Industry                                                                        | Sharchilev et al. (2018) |
|                          |                 | Age                                                                             |                          |
|                          |                 | Number of products                                                              | Te et al. (2023b)        |
|                          |                 | Presence of experienced advisors                                                |                          |
|                          | Mentions        | News article count<br>Domain-specific mentions<br>Topic modeling features (LDA) | Sharchilev et al. (2018) |
| Founders                 | –               | Number of founders                                                              | Sharchilev et al. (2018) |
|                          |                 | Demographics                                                                    |                          |
|                          |                 | Education                                                                       | Te et al. (2023b)        |
|                          |                 | Work experience                                                                 |                          |
| Investment               | Funding details | Total amount raised                                                             | Krishna et al. (2016)    |
|                          |                 | Number of rounds                                                                | Sharchilev et al. (2018) |
|                          |                 | Burn rate                                                                       |                          |
|                          |                 | Time until rounds                                                               | Te et al. (2023b)        |
|                          |                 | Capital concentration rate                                                      |                          |
|                          | Investors       | Number of investors                                                             | Wei et al. (2025)        |
|                          |                 | Investor types                                                                  |                          |
|                          |                 | Total amount invested                                                           | Te et al. (2023b)        |
|                          |                 | Average network centrality                                                      |                          |
|                          |                 | Investor success rate                                                           |                          |

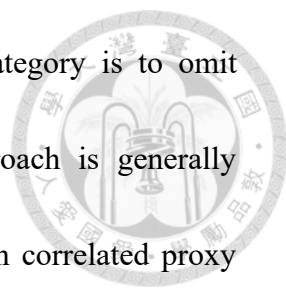
## 2.2 Existing Studies on Mitigating Unfairness in Machine Learning

Certain personal characteristics are legally recognized as impermissible bases for discrimination and are commonly referred to in the computer science literature as “protected” or “sensitive” attributes. Under the Federal Equal Credit Opportunity Act (ECOA), for example, creditors are prohibited from discriminating against credit applicants based on characteristics such as race, color, religion, national origin, sex, marital status, or age (Chen et al., 2019). In response to growing concerns that machine learning models may inadvertently learn and perpetuate biases associated with such attributes, a variety of fairness-enhancing mechanisms have been developed.

These mechanisms are typically grouped into three main categories: pre-processing, in-processing, and post-processing, each of which focuses on a specific stage of the machine learning pipeline (Binns, 2018). Pre-processing techniques aim to mitigate biases before model training by modifying the input data. In-processing methods intervene during model training to incorporate fairness considerations or modify the learning process itself. Post-processing approaches adjust the model’s outputs after training to ensure fairer decision outcomes. The following sections provide a review of representative methods within each of these categories.

### 2.2.1 Pre-Processing Methods

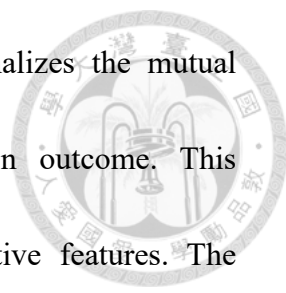
Pre-processing approaches aim to mitigate biases by modifying the training data



before model development. One intuitive strategy within this category is to omit sensitive variables from the input features. However, this approach is generally insufficient, as sensitive information can often be inferred through correlated proxy features (Pedreshi et al., 2008). To better ensure fairness in the training data, Kamiran and Calders (2012) propose a relabeling method that modifies the class labels of instances located near the decision boundary and belonging to the underrepresented group. This adjustment seeks to reduce disparities in the predicted positive rates between two groups distinguished by a sensitive attribute, while preserving the overall predictive performance of the model. Another notable method is perturbation, introduced by Feldman et al. (2015). This technique modifies the distribution of input features based on one or more binary sensitive variables, ensuring that the resulting data lacks sufficient information for classifiers to infer protected attributes. By obfuscating the link between training data and sensitive characteristics, perturbation promotes fairness by limiting the potential for indirect discrimination during subsequent model training.

### *2.2.2 In-Processing Methods*

This category of methods enforces fairness constraints during model training, either through architectural modifications or by incorporating regularization terms into the objective function. Kamishima et al. (2012) propose a fairness-aware method by

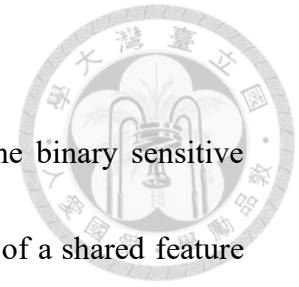


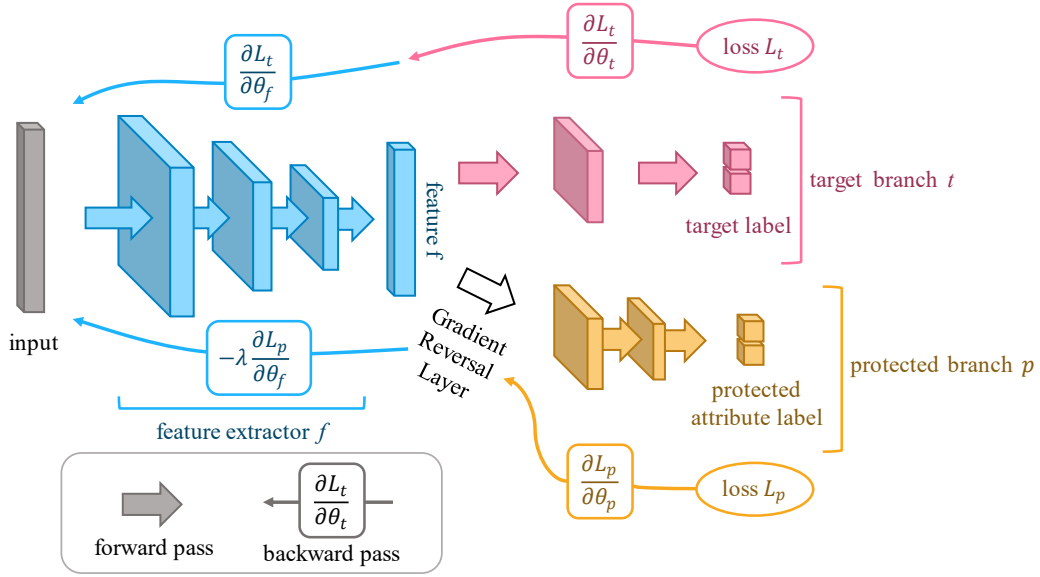
extending the loss function with a regularization term that penalizes the mutual information between the sensitive attribute and the prediction outcome. This discourages the model from making decisions based on sensitive features. The regularization term can also be designed to reflect common fairness notions such as equalized odds (Zafar et al., 2017), allowing flexible control over the model's fairness behavior. As another line of work under this category, Zemel et al. (2013) introduce Learning Fair Representations (LFR), which uses probabilistic mappings to transform raw input features into intermediate representations defined by a set of predefined prototypes. These representations aim to retain essential task-relevant information while obscuring protected group membership. By decoupling group identity from the learned features, the method seeks to ensure that predicted positive outcome probabilities are more equitably distributed across groups.

In recent work, Te et al. (2023a) apply a Gradient Reversal approach to learn representations that are invariant to sensitive attributes, particularly in the field of startup success prediction. This technique, originally introduced for domain adaptation (Ganin et al., 2016), involves optimizing two adversarial objectives simultaneously: one for the main label prediction task and another for the domain classification task. In fair classification, sensitive attributes can be analogously treated as domains, where the goal is to ensure that learned representations are predictive of the target label but

uninformative with respect to the sensitive attributes.

Accordingly, in the setting of mitigating the influence of one binary sensitive attribute, as illustrated in Figure 1, the model architecture consists of a shared feature extractor, a target prediction branch for the main task, and a protected attribute branch for predicting the sensitive attribute. The core component of this architecture is the Gradient Reversal Layer (GRL), positioned between the feature extractor and the protected branch. The GRL is a custom layer that passes inputs forward unchanged but multiplies the gradients by a negative scalar  $-\lambda$  during backpropagation. This adversarial training process forces the feature extractor to learn representations that minimize the influence of sensitive attribute while preserving task relevance, thereby promoting fairness in the final prediction outcomes.





**Figure 1: Architecture of Gradient Reversal Approach for Fair Classification**

### 2.2.3 Post-Processing Methods

Post-processing methods focus on modifying a model's predictions after it has been trained, without altering the underlying training data or model parameters. Such approaches offer practical solutions for mitigating unfairness, particularly in cases where model retraining is not feasible. For example, Hardt et al. (2016) propose a group-specific thresholding technique that adjusts the decision boundary for different demographic groups. The objective is to equalize the true positive rate and false positive rate across groups defined by a sensitive attribute, thereby promoting fairness in classification outcomes. Building on the idea of counterfactual fairness, Lohia et al. (2019) introduce a mechanism known as calibration. This method identifies individuals whose predicted outcomes would differ if their sensitive attribute were changed while holding all other features constant. The predicted labels for such individuals are then

flipped to the opposite class to reduce disparities between advantaged and disadvantaged groups.

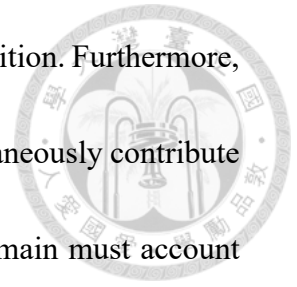


These post-processing techniques are considered as flexible as they can be combined with any classification algorithm. However, they may come at the cost of limited interpretability, as fairness is enforced externally through adjustments to model outputs that intentionally alter individual predictions rather than being learned within the model itself (Pessach & Shmueli, 2022).

### **2.3 Summary and Limitations of Existing Literature**

To illustrate the methodological landscape, Table 2 provides an overview of representative fairness-enhancing techniques. Although fairness in machine learning has received extensive attention, only a very limited number of studies have addressed fairness in startup success prediction. Most existing fairness-enhancing methods are designed for general-purpose benchmarks, typically handling only one binary sensitive attribute at a time. Moreover, they often simplify continuous variables (e.g., an individual's age, the racial composition of a community) by binarizing them to fit algorithmic constraints, thereby risking oversimplifying real-world demographic nuances. These studies also predominantly focus on individual-based fairness, assuming each prediction pertains to a single person. However, startup investment decisions are naturally team-based, where some sensitive attributes are inherently

continuous, requiring fairness notions that account for group composition. Furthermore, as discussed in Section 1.1, multiple sensitive attributes may simultaneously contribute to investor biases. Consequently, fairness considerations in this domain must account for multiple sensitive attributes of mixed data types, including both categorical and continuous variables.



**Table 2: Summary of Existing Approaches to Fairness in Machine Learning**

| Source                     | Fairness Strategy | Method                        | Datasets: Sensitive Attribute(s) (Type)                                                                               | Fairness Unit            |
|----------------------------|-------------------|-------------------------------|-----------------------------------------------------------------------------------------------------------------------|--------------------------|
| Kamiran and Calders (2012) | Pre-processing    | Relabeling                    | - Adult Income: gender (B)<br>- Communities and Crimes: race (Bc)                                                     | Individual;<br>Community |
| Feldman et al. (2015)      | Pre-processing    | Perturbation                  | - Ricci dataset: race (B)<br>- Adult Income: gender (B)<br>- German Credit: age (Bc)                                  | Individual               |
| Kamishima et al. (2012)    | In-processing     | Regularization                | - Adult Income: gender (B)                                                                                            | Individual               |
| Zemel et al. (2013)        | In-processing     | Learning Fair Representations | - Adult Income: gender (B)<br>- German Credit: age (Bc)                                                               | Individual               |
| Te et al. (2023a)          | In-processing     | Gradient Reversal             | - Crunchbase: region (B), gender (B), race (B), university (B)                                                        | Group                    |
| Hardt et al. (2016)        | Post-processing   | Thresholding                  | - FICO score dataset: race (B)                                                                                        | Individual               |
| Lohia et al. (2019)        | Post-processing   | Calibration                   | - Adult Income: gender (B) or race (B)<br>- German Credit: gender (B) or age (Bc)<br>- COMPAS: gender (B) or age (Bc) | Individual               |

B: Binary; Bc: Binary derived from continuous

**Table 3: Comparison of Fairness Methods in Handling Multiple and Mixed Type Sensitive Attributes**

| Fairness Strategy | Method                        | Multiple | Mixed |
|-------------------|-------------------------------|----------|-------|
| Pre-processing    | Blind (Omit)                  | ✓        | ✓     |
| Pre-processing    | Relabeling                    | △        | ✗     |
| Pre-processing    | Perturbation                  | △        | ✗     |
| In-processing     | Learning Fair Representations | △        | ✗     |
| In-processing     | Regularization                | ✓        | ✓     |
| In-processing     | Gradient Reversal             | ✓        | ✓     |
| Post-processing   | Thresholding                  | △        | ✗     |
| Post-processing   | Calibration                   | △        | ✗     |

△: denotes limited capability

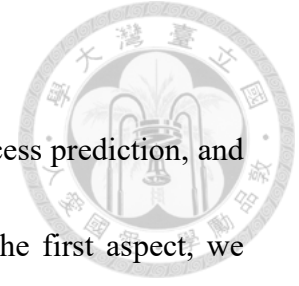
Table 3 highlights the limitations of existing methods in accommodating these complexities. Only a few approaches are capable of handling both multiple and mixed type sensitive attributes simultaneously. Most methods struggle to support multiple sensitive attributes effectively because they require discrete group boundaries and aim to obscure group membership through input manipulation, intermediate representations, or output adjustment. This is more straightforward when there is only a single binary sensitive attribute. Taking Kamiran and Calders's (2012) relabeling method as an example, they flip the negative predicted labels of female instances with prediction scores close to 0.5, in order to balance the chance of receiving positive predictions

between binary gender groups in their experiments on the Adult Income dataset. However, when extended to more realistic scenarios involving multiple sensitive attributes, this approach becomes impractical. When multiple sensitive attributes interact, it is ambiguous which intersectional groups should be considered disadvantaged, and the number of possible group combinations grows rapidly, making group definitions and fairness optimization increasingly complex.

In the field of startup success prediction, a study by Te et al. (2023a) adapts the Gradient Reversal framework to promote fairness. While their work is pioneering in applying fairness-aware learning to this domain, it simplifies continuous sensitive features by employing categorical encoding schemes. For example, team gender composition is reduced to discrete categories such as all-male, all-female, or mixed-gender, which may overlook finer-grained demographic variation.

These limitations collectively point to the need for more domain-specific fairness modeling approaches that are tailored to the unique characteristics of startups and the decision-making dynamics in venture capital. Our study builds upon prior works by developing a fairness-aware framework that explicitly accommodates multiple sensitive attributes and preserves richer representations of team composition, thus providing a more realistic and inclusive approach to fairness in startup early success prediction.

## Chapter 3 Methodology



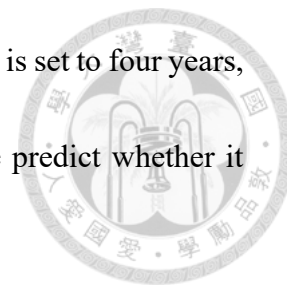
This study comprises two major aspects: (1) startup early success prediction, and (2) the design and evaluation of fairness-aware algorithms. For the first aspect, we define the criterion used to determine startup early success and present a structured overview of the predictive features, including the derivation of sensitive attributes. For the second aspect, we describe the design of our fairness-aware learning framework, which encompasses and compares multiple bias mitigation strategies, including feature exclusion (Blind), fairness-constrained regularization, and Gradient Reversal. Overall, the adapted methods aim to support equitable decision-making in VC by promoting algorithmic fairness without substantially compromising predictive performance in startup early success prediction.

### 3.1 Definition of Startup Early Success

In this study, we adopt the definition of startup early success proposed by Sharchilev et al. (2018), which conceptualizes the task as a forward-looking prediction of funding events. Specifically, we consider startups that have already achieved an early-stage funding milestone, referred to as the trigger round, and aim to predict whether they will reach the subsequent milestone, the target round, within a predefined time window.

In our implementation, angel and seed rounds are selected as trigger rounds, while

the Series A round serves as the target round. The prediction horizon is set to four years, meaning that for each startup receiving angel or seed funding, we predict whether it will obtain Series A funding within the next four years.



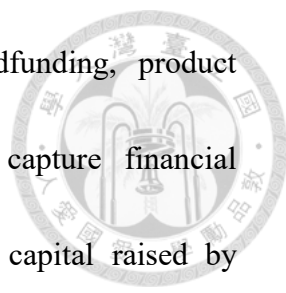
### 3.2 Predictive Features Used in Our Research

As reviewed in Section 2.1, our study incorporates three major categories of predictive features commonly employed in early-stage startup success prediction, particularly for forecasting Series A funding outcomes. These features are derived from Crunchbase, a widely used startup database that offers detailed records on company profiles, funding history, and individual founder information, making it especially suitable for analyzing fairness issues that involve sensitive personal characteristics.

- Company-Level Features: We incorporate several basic company characteristics that describe the focal startup's context and profile. These include geographic location attributes such as world region and U.S. state, as well as industry tags using categories defined by Crunchbase. In addition, we record the company age (measured in months from founding to the trigger round) and the number of advisors involved prior to the trigger event.
- Founder-Level Features: Founder-related features capture the background of the startup's founding team. Demographic attributes include gender and race, while educational background is recorded based on the subject field, degree level, and

the QS ranking of the attended university (QS Top Universities, 2024). Because the original data fields for degree and subject are provided as free-text entries, we apply a keyword-based standardization approach to unify them. This approach leverages a predefined list of degree levels and subject areas, which is constructed based on our domain expertise. Work experience is captured through job title, job type, years of experience, and serial entrepreneurship experience. A similar keyword-matching strategy used for the degree and subject fields is also applied to standardize job titles. Finally, to reflect the team composition of the focal startup founders, we aggregate these features across all founder team members, using ratios (e.g., proportion of female founders) or averages (e.g., average years of experience).

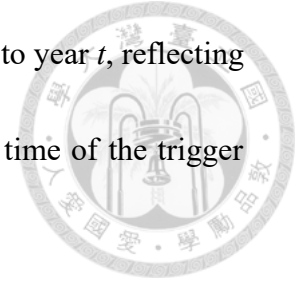
- Investment-Related Features: Investment features are further categorized into two subcategories, funding details and historical investors.
- Funding Details: This subcategory characterizes the financial development of the startup until the trigger round. These include trigger type (angel or seed), total capital raised (in USD), and number of funding rounds received. In addition to aggregated statistics across all funding types, we also record separate counts and amounts for several funding types considered likely to occur before Series A. These include pre-seed, seed, angel, convertible note,



grant, corporate round, debt financing, equity crowdfunding, product crowdfunding, and non-equity assistance. To better capture financial efficiency, we compute a burn rate by dividing total capital raised by company age (Krishna et al., 2016), and measure capital concentration rate using a prior-established formula from the literature (Wei et al., 2025), which captures the extent to which a startup's funding is dominated by a small number of investors.

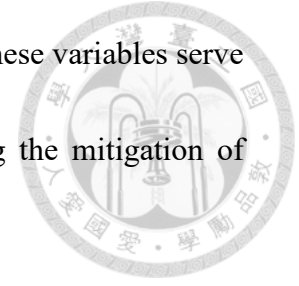
- *Historical Investors*: Features related to historical investors describe the background and network strength of the investors who have previously invested in the focal startup. We include the total and average number of historical investors, the average amount and number of investments they have made, the average number of distinct organizations they have invested in, and the count of institutional versus individual investors. Furthermore, we consider both the average and maximum historical success rate of these investors, along with their average and total network centrality. To compute network centrality, we adopt the method proposed by Wei et al. (2025), which constructs a co-investment network where two investors are linked if they co-invested in the same company in the same funding round. For a startup receiving its trigger round in year  $t$ , investor centrality is measured based on

the ten-year co-investment network built using data prior to year  $t$ , reflecting the investors' connectivity and potential influence at the time of the trigger round.



To support fairness evaluation, we identify several sensitive attributes that may influence venture funding decisions, as discussed in Section 2.1. These include both categorical and continuous variables, allowing for a more nuanced assessment of fairness. For categorical features, we designate the world region of the startup as sensitive, as home bias has been documented in VC investment decision-making (Coval & Moskowitz, 1999) and may implicitly capture socio-cultural background associated with the founders (Te et al., 2023a). For continuous features, we include the demographic composition of the founding team, specifically the proportion of female founders and the proportion of non-Caucasian founders, as sensitive attributes. Since Crunchbase does not explicitly provide racial information, we infer race from founder names using an LSTM-based text classification model. Specifically, we categorize race into six groups: European, Hispanic, East Asian, Nordic, Celtic English, and Muslim. Among these, Hispanic, East Asian, and Muslim are classified as non-Caucasian for the purpose of our fairness analysis. Although the dataset lacks direct racial labels, racial composition remains a potentially visible and discriminatory factor in venture capital investment decision-making. It is important to note that the inferred race is only an

estimated value and should not be treated as ground truth. Lastly, these variables serve as focal points in our fairness-aware modeling strategy, enabling the mitigation of potential biases.



A comprehensive list of the features employed in our methods is presented in Table A1 in the Appendix. The table summarizes variables under different categories and provides their corresponding descriptions, where sensitive features are highlighted using bold formatting. It also specifies whether each feature is classified as static or dynamic, meaning it remains constant regardless of the trigger timing, or dynamic, meaning it is time-dependent and the values may vary based on the trigger event.

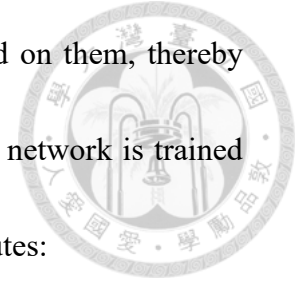
### **3.3 Fairness-Aware Modeling Approaches**

In this study, we implement and evaluate three fairness-aware methods that are inherently capable of handling multiple sensitive attributes of mixed types (categorical and continuous) as discussed in Section 2.3. These strategies aim to promote fairer prediction outcomes in early-stage startup success prediction, including one pre-processing method, the Blind method, and two in-processing methods: fairness-constrained regularization and Gradient Reversal.

#### *3.3.1 Blind Method*

The Blind method promotes fairness by excluding explicit sensitive attributes from the model input. The underlying assumption is that by removing these attributes, the

model is less likely to directly learn discriminatory patterns based on them, thereby encouraging fairer predictions. In this approach, a standard neural network is trained using a reduced feature set that omits the following sensitive attributes:



- Region: Continent-level grouping of the startup's founding location.
- Gender: Proportion of female founders.
- Race: Proportion of non-Caucasian founders.

Although conceptually straightforward, this method has notable limitations. Prior studies have shown that merely excluding sensitive attributes may not fully eliminate biases, as these attributes can often be inferred indirectly from other correlated, non-sensitive features (Pedreshi et al., 2008). Additionally, the removal of potentially predictive variables may lead to reduced model accuracy (Chen et al., 2018; Hajian & Domingo-Ferrer, 2012).

### *3.3.2 Regularization Method*

The Regularization method is a fairness-aware learning approach that incorporates fairness constraints directly into the model's loss function (Caton & Haas, 2024). Different from the Blind method, which omits sensitive attributes entirely, this strategy allows the model to access sensitive information while introducing penalties to reduce discriminatory behavior during optimization. This approach seeks to retain the predictive value of sensitive attributes while mitigating their unfair influence on model

outputs.

In particular, we add a fairness penalty term to the standard prediction loss. The penalty is defined as the sum of FairGap values associated with each protected attribute, where FairGap is a group fairness metric based on differences in mean prediction scores, formally defined in Section 4.2.1. The total loss function to be minimized is defined as:

$$\begin{aligned} L_{total} = & L_{target} + w_{region} \cdot \text{FairGap}_{region} \\ & + w_{gender} \cdot \text{FairGap}_{gender} \\ & + w_{race} \cdot \text{FairGap}_{race}, \end{aligned}$$

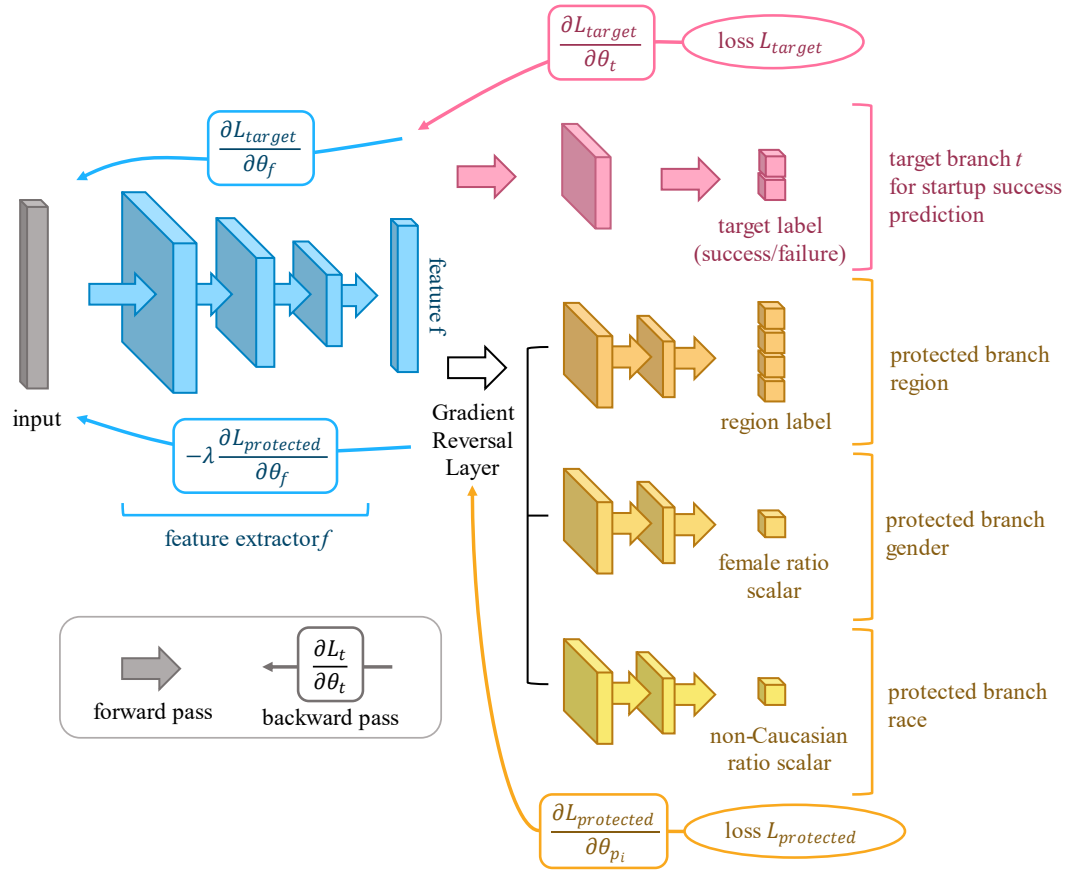
where  $L_{target}$  donates the binary cross-entropy loss for the main prediction task,  $\text{FairGap}_{attr}$  represents the fairness metric calculated for each protected attribute (region, gender, and race), and  $w_{attr}$  is the predefined weight that controls the influence of each fairness constraint on the model's training objective.

### 3.3.3 Gradient Reversal Method

Inspired by the work of Te et al. (2023a), we adopt a Gradient Reversal framework to mitigate biases in startup early success prediction. This method is based on an adversarial learning mechanism, aiming to prevent the model from capturing information related to sensitive attributes, including signals that may be hidden in correlated but non-sensitive features, during the process of representation learning.

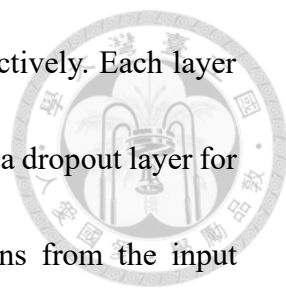
The model architecture consists of a feature extractor shared by two types of branches: a target prediction branch for predicting startup success, and multiple

protected attribute branches for predicting sensitive attributes such as region, gender, and race. To promote fairness, a Gradient Reversal Layer (GRL) is inserted between the feature extractor and each protected branch. During backpropagation, the GRL multiplies the gradient flowing into the feature extractor by a negative scalar  $-\lambda$ , thereby encouraging the model to learn representations that are predictive of the target but invariant to the sensitive attributes. Figure 2 illustrates the overall framework of our Gradient Reversal-based model.



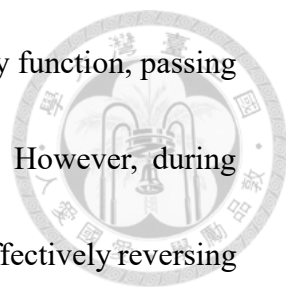
**Figure 2: Architecture of Our Gradient Reversal Approach for Fair Startup Success Prediction**

The feature extractor  $f$  is a feedforward neural network consisting of three fully



connected layers with hidden dimensions of 128, 64, and 32, respectively. Each layer is followed by batch normalization, a ReLU activation function, and a dropout layer for regularization. This component aims to learn latent representations from the input features and serves as a shared input to both the target and protected branches. The target branch  $t$  is a simple classifier for predicting startup success. It consists of a single linear layer that outputs a scalar logit, which is then passed through a sigmoid activation to obtain the predicted probability for binary classification. This branch is trained to maximize predictive performance based on the learned latent representations using a binary cross-entropy loss function  $L_{target}$ , which measures the discrepancy between the predicted probabilities and the ground truth labels. A set of protected branches, each denoted as  $p_{attr}$ , is built to predict certain sensitive attributes from the shared representation. Each protected branch is composed of a two-layer feedforward neural network with hidden dimensions of 16 and 8, where both layers are followed by a ReLU activation function. The output dimension of the second layer corresponds to the number of categories for categorical attributes, or one scalar for continuous ones. The loss function  $L_{p_{attr}}$  for each protected attribute is defined as cross-entropy for categorical attributes and mean squared error (MSE) for continuous attributes.

To enable adversarial training against sensitive attribute prediction, a Gradient Reversal Layer (GRL) is inserted between the feature extractor and each protected



branch. During forward propagation, the GRL behaves as an identity function, passing the learned representations unchanged to the protected branches. However, during backpropagation, it multiplies the gradient by a negative scalar  $-\lambda$ , effectively reversing the direction of the gradient flow. This discourages the feature extractor from encoding information predictive of sensitive attributes, thereby promoting fairer representations that are invariant to these factors. The GRL itself contains no learnable parameters. Finally, to regulate the strength of the adversarial signal, we adopt a scheduled  $\lambda$  strategy following the approach proposed by Ganin et al. (2016). The scalar  $\lambda$  is gradually increased throughout training according to a predefined schedule based on the normalized training progress  $p \in [0, 1]$ :

$$\lambda(p) = \lambda_{\max} \cdot \left( \frac{2}{1 + \exp(-10p)} - 1 \right)$$

This schedule starts with a small  $\lambda$ , allowing the model to focus on learning predictive features in the early steps. As training progresses,  $\lambda$  increases and stabilizes near a maximum value  $\lambda_{\max}$ , progressively enforcing stronger fairness constraints through adversarial pressure.

To further enhance training stability and ensure effective learning, we adopt a two-stage optimization procedure using three separate optimizers: one for the feature extractor (Optimizer F), one for the target branch (Optimizer T), and one for the protected branches (Optimizer P). This training strategy is designed to decouple the

competing objectives of accurate prediction of the main task and fairness enhancement.

In the first stage, we focus on the target prediction task. The binary cross-entropy loss

$L_{target}$  is computed from the output of the target branch then backpropagated to update

the parameters of both the feature extractor and the target branch using Optimizer F and

Optimizer T, respectively. In the second stage, we shift to the fairness objective. For

each protected branch of sensitive attribute  $attr$ , the loss  $L_{p_{attr}}$  is computed.

Subsequently, each loss is scaled by a predefined weight  $w_{attr}$  that controls its relative

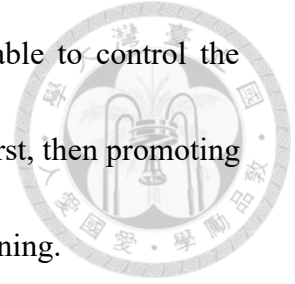
influence, and the total protected loss  $L_{protected}$  is computed as a weighted sum:

$$L_{protected} = \sum_{attr} w_{attr} \cdot L_{p_{attr}}$$

The GRL is then applied to reverse the gradients of the total loss from the protected branches before they are backpropagated to the feature extractor. Finally, Optimizer F and Optimizer P are used to update the feature extractor and protected branches, respectively.

This two-stage training scheme helps avoid gradient conflicts between the competing objectives. If a joint loss function combining the target prediction task and the fairness objectives is used to update the feature extractor simultaneously, it will receive opposing gradient signals, one from the target task that encourages preserving predictive information, and another from the adversarial branches that encourages removing sensitive attribute signals. These conflicting directions can lead to unstable

learning. By separating the updates into distinct stages, we are able to control the training signals more precisely by reinforcing predictive capacity first, then promoting fairness, thus enabling more stable and balanced representation learning.



## Chapter 4 Experiments



### 4.1 Data Collection

We construct our dataset based on a snapshot of Crunchbase as of April 1, 2025, focusing on organizations whose primary role is classified as “company” and having a “Software” category tag to ensure industry relevance. To align with our four-year prediction horizon, we retain only those companies founded between January 1, 2011, and April 1, 2021, resulting in an initial pool of 312,996 candidate startups. Several filtering criteria are then applied to refine the dataset. First, the company must have received a trigger round (i.e., seed or angel round), yielding 41,194 companies. Second, companies that received over \$10 million USD in funding within one year of founding or took more than four years to raise a trigger round are excluded, as the former are often spinoffs of large corporates and the latter, which managed to sustain operations for an extended period without raising early funding, may differ from the typical early-stage startups. This results in 40,880 remaining firms. Third, companies lacking valid founder records are removed, leaving 11,274 candidates. After excluding firms with missing historical investor information, the final dataset comprises 13,592 funding instances from 9,206 unique startups, with 33.98% of them labeled as positive cases (i.e., for reaching Series A within four years). Table A2 in the Appendix provides the descriptive statistics of our dataset.

Additional preprocessing steps are applied to the data. Continuous variables exhibiting strong right skewness (i.e., skewness  $> 1$ ) are log-transformed to mitigate long-tail effects and stabilize variance. To ensure consistent scaling across features and facilitate model convergence, all continuous variables not originally bounded within  $[0, 1]$  are normalized to the interval of  $[0, 1]$ . For categorical features, except for sensitive attributes, one-hot encoding is applied to ensure compatibility with the neural network input structure.

## 4.2 Baseline and Model Settings

To validate the effectiveness of the fairness-aware methods investigated, we compare them against a baseline model referred to as Standard NN. This model is a standard feedforward neural network trained without incorporating any fairness constraints. We aim to assess improvements in fairness metrics, as well as potential trade-off in predictive effectiveness, by comparing the fairness-aware methods against this baseline setting.

**Table 4: Model Settings of Baseline and Fairness-Aware Approaches**

| Model             | Fairness Method        | Input Features            | Protected Branches | $\lambda_{\max}$ |
|-------------------|------------------------|---------------------------|--------------------|------------------|
| Standard NN       | ✗                      | Full (incl. sensitive)    | ✗                  | ✗                |
| Blind             | Omit                   | Partial (excl. sensitive) | ✗                  | ✗                |
| Regularization    | FairGap regularization | Full (incl. sensitive)    | ✗                  | ✗                |
| Gradient Reversal | Gradient reversal      | Partial (excl. sensitive) | ✓                  | 3.2              |

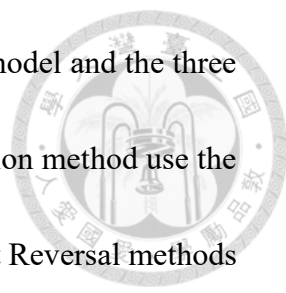


Table 4 summarizes the model configurations of the baseline model and the three fairness-aware methods. Both the Standard NN and the Regularization method use the complete feature set as the input. In contrast, the Blind and Gradient Reversal methods use a version of the dataset that excludes the three sensitive attributes from the input features. Among the four methods, only the Gradient Reversal method requires the construction of protected branches and the specification of a  $\lambda_{\max}$  parameter to regulate the strength of adversarial learning. In our experiments,  $\lambda_{\max}$  is set to 3.2. All methods share the same architecture for the feature extractor and the target prediction branch. Specifically, the feature extractor is implemented as three fully connected layers with hidden dimensions of 128, 64, and 32, each followed by batch normalization, a ReLU activation, and dropout (rate = 0.2). The target branch is a linear layer mapping the final hidden representation to a scalar logit, trained using binary cross-entropy loss. All methods are trained using the Adam optimizer with a learning rate of 1e-3, batch size of 2048, and a total of 100 epochs. The input dimension is 153 when using the complete dataset, and 150 when sensitive attributes are excluded. For fairness-aware methods that handle multiple sensitive attributes, fixed loss weights are assigned to each attribute: 0.3 for region, 0.8 for gender, and 0.8 for race.

## 4.3 Evaluation Design

### 4.3.1 Fairness Evaluation Metrics



Traditional fairness metrics are typically designed to compare disparities between binary groups, such as privileged and unprivileged populations, making them less suitable for continuous or multi-class sensitive attributes. Therefore, inspired by FairQuant (Grari et al., 2019), we propose a modified metric, FairGap, which offers greater flexibility in handling both continuous and categorical sensitive features. The original FairQuant measures disparities by dividing samples into equal-sized bins based on quantiles of a continuous sensitive attribute, then calculating the average absolute difference between each bin's prediction mean and the global average. However, this may result in multiple bins containing similar or even identical values of the sensitive attribute, allowing majority groups with comparable attribute levels to dominate the fairness metric and potentially obscure disparities affecting minority groups. Consequently, instead of using quantile-based binning to enforce equal sample sizes, we group samples based on evenly spaced value intervals of the continuous sensitive attribute, where each interval represents a distinct group with a certain attribute range. This allows for more semantically meaningful groupings. Furthermore, rather than comparing each group to a global average, we compute pairwise absolute differences across all group combinations, which better aligns with the traditional notion of group

fairness as the disparity between any two subpopulations. Ultimately, this pairwise formulation naturally extends to multi-class sensitive attributes, thereby supporting a unified evaluation framework across attribute types. The detailed definition of FairGap and its extension to multi-class settings will be discussed later in this section.

In our experiments, for continuous sensitive attributes (e.g., female founder ratio), we discretize samples into four groups based on their value intervals (i.e.,  $[0, 0.25)$ ,  $[0.25, 0.5)$ ,  $[0.5, 0.75)$ , and  $[0.75, 1]$ ). We then compute the pairwise differences of a relevant fairness metric (e.g., positive prediction rate for demographic parity) across all group pairs and average the results to obtain the final FairGap score:

$$\text{FairGap} = \frac{2}{k(k-1)} \sum_{i < j} |s_i - s_j|$$

where  $k$  is the number of bins;  $s_i$  and  $s_j$  denote the fairness score for group  $i$  and  $j$ , respectively.

**Table 5: FairGap Computation Example Using Gender as Sensitive Attribute**

| Group: Female Founder Ratio | Positive Prediction Rate | Absolute Pairwise Differences         | FairGap Computation                                                                                                                                         |
|-----------------------------|--------------------------|---------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------|
| A: 0–25%                    | 0.60                     | A–B: 0.10,<br>A–C: 0.15,<br>A–D: 0.05 | <b>Sum of all absolute pairwise differences:</b><br>$0.10 + 0.15 + 0.05 + 0.05 + 0.05 + 0.10 = 0.50$<br><br><b>FairGap score:</b><br>$0.50 \div 6 = 0.0833$ |
| B: 25–50%                   | 0.50                     | B–C: 0.05,<br>B–D: 0.05               |                                                                                                                                                             |
| C: 50–75%                   | 0.45                     | C–D: 0.10                             |                                                                                                                                                             |
| D: 75–100%                  | 0.55                     | –                                     |                                                                                                                                                             |

To illustrate how FairGap is computed, Table 5 shows a toy example using four groups based on female founder ratio and their corresponding positive prediction rates. The pairwise absolute differences in positive prediction rates are 0.10, 0.15, 0.05, 0.05, 0.05, and 0.10, summing to 0.50. The FairGap is calculated as the average of six group pairs, which in this case is approximately equal to 0.0833.

This formulation allows FairGap to flexibly incorporate various fairness score definitions. In our case, we focus on positive prediction rate and report the corresponding FairGap scores of three sensitive attributes in our results. For categorical attributes such as region, we treat each category as a distinct group and apply the same pairwise comparison procedure. A lower FairGap score indicates smaller disparities and

hence greater fairness with respect to the given sensitive attribute.

In addition to group-level fairness, we also evaluate individual-level fairness using the consistency metric, a commonly adopted measure in algorithmic fairness research.

Consistency assesses whether individuals with similar attributes receive similar model predictions (Dwork et al., 2012). Specifically, for each instance in the dataset, the predicted outcome is compared against those of its  $k$  nearest neighbors in the input feature space:

$$\text{Consistency} = 1 - \frac{1}{n} \sum_{i=1}^n \left| \hat{y}_i - \frac{1}{k} \sum_{j \in N_k(i)} \hat{y}_j \right|$$

In this equation,  $n$  denotes the total number of instances,  $\hat{y}_i$  represents the predicted label for instance  $i$ ,  $N_k(i)$  refers to the set of  $k$  nearest neighbors of instance  $i$ , and  $\hat{y}_j$  are the predicted labels for these neighbors. A higher consistency score indicates greater individual-level fairness, with a maximum score of 1 implying perfect consistency where every instance receives a prediction identical to those of its nearest neighbors. In our experiments, we set  $k$  to 5 following common practice in the literature.

To illustrate how consistency is calculated, consider a simple example. Suppose an instance receives a predicted label of 1, and its five nearest neighbors have predicted labels of 0, 0, 0, 1, and 1. Among these neighbors, only two share the same prediction as the instance itself. Therefore, the individual consistency for this sample is  $2/5 = 0.4$ .

By computing this value for every sample in the testing set and then averaging across

all samples, we obtain the final consistency score, which reflects the overall individual-level fairness of the model.



#### *4.3.2 Evaluation Procedure and Performance Metrics*

To evaluate the effectiveness of our methods on the primary prediction task, that is, whether a startup secures Series A funding within four years after receiving angel or seed financing, we adopt several standard binary classification metrics. Specifically, we report the following indicators:

- Accuracy
- AUC (area under curve)
- Precision, Recall, and F1-score for the positive class (i.e., successful cases)
- Precision, Recall, and F1-score for the negative class (i.e., unsuccessful cases)

These metrics capture both general performance and class-specific behavior, which is particularly important given the class imbalance in our dataset.

All methods are evaluated using repeated 10-fold cross-validation, conducted 30 times with different random splits. Metrics are computed independently for each fold and repetition, and the final results are reported as averages across all evaluations. In addition to predictive performance, we report the fairness metrics introduced in the previous subsection. A fixed decision threshold of 0.5 is used across all methods to determine binary predictions from output probabilities.



Table 6: Evaluation Results of Baseline and Fairness-Aware Methods

| Method         | Acc.          | AUC           | Precision<br>_pos | Recall<br>_pos | F1_pos        | Precision<br>_neg | Recall<br>_neg | F1_neg        | Consis-<br>tency | FairGap<br>_region | FairGap<br>_gender | FairGap<br>_race |
|----------------|---------------|---------------|-------------------|----------------|---------------|-------------------|----------------|---------------|------------------|--------------------|--------------------|------------------|
| Standard<br>NN | <b>77.11%</b> | <b>82.32%</b> | <u>66.64%</u>     | <b>66.16%</b>  | <b>66.20%</b> | <b>82.69%</b>     | 82.74%         | <u>82.65%</u> | 0.6304           | 0.0987             | 0.1229             | 0.1105           |
| Blind          | 76.57%        | <u>81.75%</u> | 65.76%            | <u>65.69%</u>  | <u>65.53%</u> | <u>82.39%</u>     | 82.17%         | 82.22%        | 0.6298           | 0.082              | 0.1081             | 0.1121           |
|                | (-0.54%)      | (-0.57%)      | (-0.88%)          | (-0.47%)       | (-0.67%)      | (-0.30%)          | (-0.57%)       | (-0.43%)      | (-0.0006)        | (-0.0167)          | (-0.0148)          | (+0.0016)        |
| GR             | <u>76.60%</u> | 81.23%        | <b>67.73%</b>     | 60.19%         | 63.50%        | 80.67%            | <b>85.04%</b>  | <b>82.74%</b> | <u>0.6491</u>    | <b>0.0708</b>      | <u>0.1006</u>      | <u>0.1066</u>    |
|                | (-0.51%)      | (-1.09%)      | (+1.09%)          | (-5.97%)       | (-2.70%)      | (-2.02%)          | (+2.30%)       | (+0.09%)      | (+0.0187)        | (-0.0279)          | (-0.0223)          | (-0.0039)        |
| Reg.           | 75.05%        | 79.48%        | 65.41%            | 57.34%         | 60.81%        | 79.41%            | <u>84.17%</u>  | 81.65%        | <b>0.6699</b>    | <u>0.0778</u>      | <b>0.0714</b>      | <b>0.0501</b>    |
|                | (-2.06%)      | (-2.84%)      | (-1.23%)          | (-8.82%)       | (-5.39%)      | (-3.28%)          | (+1.43%)       | (-1.00%)      | (+0.0395)        | (-0.0209)          | (-0.0515)          | (-0.0604)        |

**Note:** The **bold values** indicate the best performance for each metric, while underlined values denote the second best. For metrics from accuracy to consistency, higher values indicate better performance; for FairGap, lower values reflect greater fairness. Values in parentheses represent the difference from the Standard NN model for the corresponding metric.

## 4.4 Evaluation Results

### 4.4.1 Fairness-aware Methods



Table 6 compares the baseline and fairness-aware methods across the predictive effectiveness and the fairness metrics. Both the Blind and Gradient Reversal (GR) methods result in relatively minor reductions in accuracy and AUC compared to the baseline, suggesting a smaller trade-off in predictive performance. In contrast, the Regularization (Reg.) method shows a more pronounced decline, indicating a larger compromise in model utility. Differences in F1-score for the positive class are more evident. Blind shows only a minor drop of 0.67%, whereas Gradient Reversal and Regularization yield larger declines of 2.7% and 5.39%, respectively. Further examining the precision and recall for the positive class, we observe that compared to the baseline, the Gradient Reversal method improves precision by 1.09%, indicating a better ability to correctly identify successful startups. This could help investors avoid incorrect positive predictions and reduce the risk of misinformed investment decisions. On the other hand, all fairness-aware methods show a decline in recall for the positive class, particularly Gradient Reversal and Regularization, suggesting a reduced capacity to detect promising startups, potentially overlooking viable investment opportunities. Additionally, improvements in recall for the negative class under Gradient Reversal and Regularization indicate a stronger tendency to make negative predictions. This shift

suggests these methods achieve fairness partly by avoiding overestimation of success for certain groups. We further investigate this pattern later in the next section by analyzing the group-wise positive prediction rates across different sensitive attributes.

With regard to the individual-level fairness metric consistency, although the Blind method experiences the smallest effectiveness loss among fairness-aware methods, it fails to achieve improved fairness, as its consistency score declines compared to the baseline. By contrast, the Gradient Reversal and Regularization methods improve consistency by 0.0187 and 0.0395, respectively, indicating meaningful advancements in individual fairness. Regarding the FairGap metric, the Blind method shows limited improvement and even results in a higher FairGap for the race attribute, indicating worsened fairness. In comparison, both the Gradient Reversal and Regularization methods consistently reduce FairGap across all three sensitive attributes. The Gradient Reversal method brings the largest reduction in region FairGap, indicating the greatest improvement in regional fairness. Notably, the Regularization method achieves larger reductions in FairGap for gender and race, where the other methods show smaller improvements.

One plausible explanation for these differences in group-level unfairness mitigation effectiveness lies in the correlation between sensitive and non-sensitive features. To investigate this, we compute the Pearson correlation coefficients between



sensitive features and non-sensitive ones. The results show relatively strong correlations ( $|r| > 0.2$ ) between region and several non-sensitive features, including the network centrality of historical investors, the average number of past investments made by historical investors, and the state variable. In particular, the correlation between the North America region and the state variable reaches as high as -0.928, suggesting that regional information is likely embedded in other features. As for gender and race, they show weaker associations with non-sensitive features. The highest observed correlation for gender is 0.112 (with the Community and Lifestyle industry tag), and for race is 0.073 (with the state variable). These findings suggest that even after excluding sensitive attributes, some information can still be indirectly inferred from remaining features, which limits the effectiveness of the Blind method. On the other hand, the Gradient Reversal method improves fairness by minimizing the presence of sensitive information in the learned representations, when such information is embedded within non-sensitive features. This explains its relatively stronger performance on regional fairness and weaker performance on gender and race, whose signals are less likely to be captured in the rest of the data. Instead, the Regularization method is unaffected by the underlying correlations among variables. By directly penalizing outcome disparities measured by FairGap for a given sensitive attribute, it can reduce group-level unfairness straightforwardly, regardless of whether sensitive information is embedded in the

features.

Taken together, the results reveal a trade-off: the methods that achieve higher levels of fairness tend to bring on greater reductions in predictive effectiveness. The Regularization method, which shows the most comprehensive improvement in fairness metrics, experiences a 5.39% decrease in F1 score for the positive class and a 1% decrease for the negative class. These performance drops are viewed as an acceptable trade-off in pursuit of fairness, as the fundamental objective of fairness-aware algorithms is not merely to maximize classification performance, but to generate more equitable and socially responsible outcomes. In the long term, such methods can help prevent the reinforcement of historical biases and lead to more favorable results for all stakeholders involved.

#### *4.4.2 Group-wise Disparities across Sensitive Attributes*

We further examine differences in positive prediction rates across groups defined by sensitive attributes to assess whether the model without fairness interventions tends to systematically favor or disfavor certain types of startups, as well as to illustrate in greater detail how fairness-aware methods mitigate such disparities. For the region attribute, startups are categorized into four groups: North America, Europe, East Asia, and Other. Table 7 presents the positive prediction rates across these groups for all methods. In the baseline model, startups from Europe received the lowest rate of

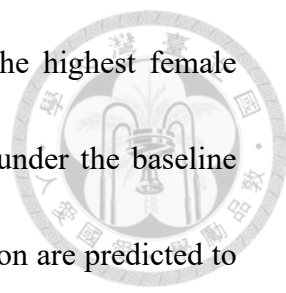
positive predictions, with only 26.10% predicted to succeed, whereas 39.64% of East Asian startups were predicted as successful. This indicates substantial disparities in receiving positive predictions across regions, with a gap of 13.54% between these two groups. Among the more effective fairness-aware methods, namely Gradient Reversal and Regularization, we observe distinct mitigation patterns. The Gradient Reversal method does not increase the positive prediction rate of the underrepresented group; instead, it reduces the rate for high-scoring groups, such as East Asia and North America, thereby narrowing the prediction gap. Regularization exhibits a similar pattern, but with a slight increase in the positive rate for the European group, suggesting a more balanced treatment by giving relatively more opportunities to underrepresented startups.

**Table 7: Positive Prediction Rates by Region Across Methods**

| Method             | Region                             |                          |                             |                         |
|--------------------|------------------------------------|--------------------------|-----------------------------|-------------------------|
|                    | Positive Rate<br>_North<br>America | Positive Rate<br>_Europe | Positive Rate<br>_East Asia | Positive Rate<br>_Other |
| <b>Standard NN</b> | 38.38%                             | <b>26.10%</b>            | <b>39.64%</b>               | 28.51%                  |
| <b>Blind</b>       | 38.76%                             | <b>27.13%</b>            | <b>28.55%</b>               | 28.72%                  |
| <b>GR</b>          | 34.02%                             | 24.89%                   | <b>25.49%</b>               | 26.03%                  |
| <b>Reg.</b>        | 33.01%                             | <b>26.78%</b>            | <b>21.54%</b>               | 25.03%                  |

**Note:** Red values indicate the lowest positive prediction rate among groups in the baseline model, while green values indicate the highest. Bold values highlight cases where a fairness method either increases the lowest rate or decreases the highest rate, thereby helping to mitigate inter-group disparity.

To reflect varying levels of gender composition of startup teams, we divide the test samples into four groups according to predefined value intervals of the female founder ratio: <25%, 25–50%, 50–75%, and >75%, which aligns with the group settings used



in the FairGap calculation. As shown in Table 8, startups with the highest female representation (>75%) receive the lowest positive prediction rate under the baseline model, at only 20.68%, while those with 25–50% female composition are predicted to succeed most frequently, with a positive prediction rate of 43%, revealing a disparity of 22.32% between these two groups. All fairness-aware methods are able to increase the positive prediction rate for the most underrepresented group. It is noteworthy that both Gradient Reversal and Regularization approaches also reduce the prediction rate for the favored group, thus narrowing the disparity. Among these, the Regularization-based method achieves the greatest reduction, lowering the group difference to 8.46%, indicating its effectiveness in mitigating gender-related prediction imbalances.

**Table 8: Positive Prediction Rates by Gender Composition Across Methods**

| Method             | Gender (female founder ratio) |                           |                           |                          |
|--------------------|-------------------------------|---------------------------|---------------------------|--------------------------|
|                    | Positive Rate<br>_< 25%       | Positive Rate<br>_25%-50% | Positive Rate<br>_50%-75% | Positive Rate<br>_> 75 % |
| <b>Standard NN</b> | 34.59%                        | <b>43.00%</b>             | 29.91%                    | <b>20.68%</b>            |
| <b>Blind</b>       | 34.07%                        | 45.53%                    | 33.99%                    | <b>25.74%</b>            |
| <b>GR</b>          | 30.24%                        | <b>41.61%</b>             | 29.98%                    | <b>23.40%</b>            |
| <b>Reg.</b>        | 30.63%                        | <b>31.37%</b>             | 26.00%                    | <b>22.91%</b>            |

**Note:** Annotations follow the same conventions as in Table 7.

We apply a similar grouping strategy as used for the female founder ratio to divide the test samples based on the proportion of non-Caucasian founders. The resulting prediction disparities across groups are presented in Table 9. Startups with 25–50% non-Caucasian founders received the highest positive prediction rate of 50%, whereas

those with more than 75% non-Caucasian representation received the lowest, only 30.93%. These indicates substantial disparities across demographic compositions, as the difference in positive prediction rates between the aforementioned two groups reaches 19.07%. Both the Gradient Reversal and Regularization methods address these disparities primarily by reducing the positive prediction rates of the favored groups, thereby narrowing the largest inter-group differences to 6.86%.

**Table 9: Positive Prediction Rates by Racial Composition Across Methods**

| Method             | Race (non-Caucasian founder ratio) |                           |                           |                          |
|--------------------|------------------------------------|---------------------------|---------------------------|--------------------------|
|                    | Positive Rate<br>_< 25%            | Positive Rate<br>_25%-50% | Positive Rate<br>_50%-75% | Positive Rate<br>_> 75 % |
| <b>Standard NN</b> | 31.88%                             | <b>50.00%</b>             | 38.78%                    | <b>30.93%</b>            |
| <b>Blind</b>       | 32.10%                             | 50.18%                    | 39.45%                    | 30.80%                   |
| <b>GR</b>          | 28.55%                             | <b>45.76%</b>             | 34.98%                    | 27.15%                   |
| <b>Reg.</b>        | 30.41%                             | <b>33.76%</b>             | 30.48%                    | 26.90%                   |

**Note:** Annotations follow the same conventions as in Table 7.

The above results show that these fairness-aware methods ensure fairness mainly by reducing the overestimation of favored groups and addressing the underestimation of unfavored groups, which may arise from discriminatory biases embedded in the data. In this sense, the benefits of such methods lie not only in helping investors avoid biased decision-making and the risk of resource misallocation, but also in facilitating fairer access to capital for high-potential ventures that might otherwise be overlooked due to systemic inequities, ultimately making the entrepreneurial ecosystem more inclusive and equitable.

#### 4.4.3 Attribution of Prediction Unfairness to Sensitive Attributes

In the following experiments, we aim to identify which sensitive attribute contributes most to prediction unfairness. To do so, we conduct controlled experiments where only one sensitive attribute is included in the dataset at a time, while the other two are excluded. This design allows us to isolate the potential impact of each sensitive attribute on model biases.

As shown in Table 10, when using the Standard NN model, the consistency scores remain similar across all three single-attribute settings, suggesting that each sensitive attribute impacts comparably to individual fairness. However, when only the sensitive attribute gender is included, we observe the highest FairGap, indicating that gender is a significant source of group-level unfairness in the baseline model.

**Table 10: Impact of Isolated Sensitive Attributes on Fairness Metrics**

| Method      | Region        |                | Gender        |                | Race          |               |
|-------------|---------------|----------------|---------------|----------------|---------------|---------------|
|             | Consistency   | FairGap_region | Consistency   | FairGap_gender | Consistency   | FairGap_race  |
| Standard NN | <u>0.6353</u> | 0.0956         | <u>0.6339</u> | 0.1228         | 0.6319        | 0.1136        |
| Blind       | 0.6317        | 0.0824         | 0.6331        | 0.1088         | <u>0.6343</u> | 0.1137        |
| GR          | <b>0.6459</b> | <b>0.0701</b>  | 0.6298        | <u>0.1037</u>  | 0.6293        | <u>0.1073</u> |
| Reg.        | 0.6309        | <u>0.0718</u>  | <b>0.6588</b> | <b>0.0782</b>  | <b>0.6544</b> | <b>0.0975</b> |

**Note:** The bold values indicate the best performance for each metric, while underlined values denote the second best. For consistency, higher values indicate better performance; for FairGap, lower values reflect greater fairness.

Furthermore, applying fairness interventions to address the bias associated with each attribute also yields meaningful improvements. Specifically, the Gradient Reversal method is more effective in improving fairness related to region, while the

Regularization method demonstrates greater improvements in gender and race fairness.

These findings align with the results discussed in Section 4.4.1.

#### 4.4.4 Impact of Fairness Interventions on A Single Sensitive Attribute

We conduct additional experiments to investigate whether applying fairness interventions to only one sensitive attribute, while all three attributes remain present in the dataset, would jointly harm or benefit fairness with respect to the others. As shown in Table 11, when fairness is enforced solely on region, the Gradient Reversal method improves all fairness metrics, including both consistency and group fairness across the other two sensitive attributes. In contrast, both the Blind and Regularization methods yield improvements only for region, while degrading fairness for all other metrics.

**Table 11: Impact of Fairness Interventions on Region Only**

| Method             | Consistency | FairGap_region | FairGap_gender | FairGap_race |
|--------------------|-------------|----------------|----------------|--------------|
| <b>Standard NN</b> | 63.04%      | 0.0987         | 0.1229         | 0.1105       |
| <b>Blind</b>       | 62.71% (↓)  | 0.0809 (↓)     | 0.1242 (↑)     | 0.1110 (↑)   |
| <b>GR</b>          | 64.08% (↑)  | 0.0718 (↓)     | 0.1158 (↓)     | 0.1056 (↓)   |
| <b>Reg.</b>        | 62.74% (↓)  | 0.0710 (↓)     | 0.1295 (↑)     | 0.1178 (↑)   |

**Note:** Arrows indicate the direction of change relative to the Standard NN model: (↑) means an increase, (↓) means a decrease. For consistency, higher values indicate better performance (↑); for FairGap, lower values indicate better fairness (↓).

Table 12 presents the experimental results when a fairness-aware method is applied only to gender. Again, the Gradient Reversal method demonstrates the ability to reduce FairGap across all three sensitive attributes, though it does not lead to improvements in consistency. The Regularization method, on the other hand, improves

fairness metrics related to gender and race, but not region. A similar pattern is observed in Table 13, where fairness is applied exclusively to race. Gradient Reversal improves all fairness metrics, while Regularization shows the same pattern that enhancing fairness in race and gender, but fails to improve regional fairness.

**Table 12: Impact of Fairness Interventions on Gender Only**

| Method             | Consistency | FairGap_region | FairGap_gender | FairGap_race |
|--------------------|-------------|----------------|----------------|--------------|
| <b>Standard NN</b> | 63.04%      | 0.0987         | 0.1229         | 0.1105       |
| <b>Blind</b>       | 63.04% (-)  | 0.0973 (↓)     | 0.1083 (↓)     | 0.1128 (↑)   |
| <b>GR</b>          | 62.79% (↓)  | 0.0934 (↓)     | 0.1062 (↓)     | 0.1081 (↓)   |
| <b>Reg.</b>        | 65.06% (↑)  | 0.1025 (↑)     | 0.0741 (↓)     | 0.0992 (↓)   |

**Note:** Annotations follow the same conventions as in Table 11.

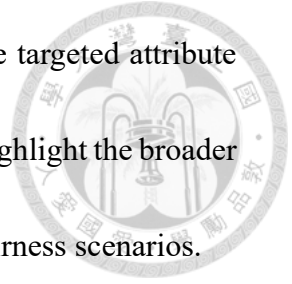
**Table 13: Impact of Fairness Interventions on Race Only**

| Method             | Consistency | FairGap_region | FairGap_gender | FairGap_race |
|--------------------|-------------|----------------|----------------|--------------|
| <b>Standard NN</b> | 63.04%      | 0.0987         | 0.1229         | 0.1105       |
| <b>Blind</b>       | 63.02% (↑)  | 0.0949 (↓)     | 0.1233 (↑)     | 0.1089 (↓)   |
| <b>GR</b>          | 63.34% (↑)  | 0.0936 (↓)     | 0.1187 (↓)     | 0.1039 (↓)   |
| <b>Reg.</b>        | 65.10% (↑)  | 0.1021 (↑)     | 0.0727 (↓)     | 0.0974 (↓)   |

**Note:** Annotations follow the same conventions as in Table 11.

These results highlight that the Gradient Reversal method can promote fairness across multiple attributes even when explicitly targeting only one. This spillover effect is likely attributable to the overlap or correlation among sensitive attributes. By mitigating the sensitive information associated with the targeted attribute in the learned representations, the Gradient Reversal method can also suppress correlated information from other attributes, thereby promoting fairness beyond the intended target. In contrast,

the Regularization method is effective at improving fairness for the targeted attribute but offers limited benefits for untargeted attributes. These findings highlight the broader potential of gradient reversal-based approaches in multi-attribute fairness scenarios.



## Chapter 5 Conclusion



### 5.1 Conclusion

To address the issue of fairness in startup early success prediction, we implement and compare three fairness-aware methods that natively support multiple sensitive attributes of mixed data types: feature-blind learning, regularization-based training, and gradient reversal. When compared to a standard neural network model, the feature-blind method fails to improve fairness effectively. The Gradient Reversal method improves fairness across all sensitive attributes, with notable gains in individual fairness and regional fairness, and moderate improvements in gender and racial fairness. Besides enhancing consistency, the Regularization method demonstrates a more uniform reduction of FairGap across all sensitive attributes. Finally, although pursuing fairness comes with a performance trade-off, our results demonstrate that this trade-off is modest, suggesting that greater fairness can be achieved without severely compromising predictive performance.

Building on this, our experiments further highlight that gender is the most significant contributor to group-level prediction unfairness. Specifically, startups with over 75% female representation among founders are the most disadvantaged under the baseline model, receiving the lowest rates of positive predictions. All three fairness-aware methods are able to improve positive outcomes for this underprivileged group,

underscoring the effectiveness of these methods to mitigate gender-based bias.

Lastly, we examine the influence of applying fairness interventions to individual sensitive attributes and find important differences in method behavior. The Gradient Reversal method demonstrates a unique advantage in its ability to simultaneously enhance fairness across multiple sensitive attributes, even when only one attribute was explicitly targeted. In contrast, the Regularization method primarily improves fairness for the targeted sensitive attribute, with limited impact on others.

In conclusion, the Regularization and Gradient Reversal methods offer distinct advantages. The Regularization method introduces explicit penalty terms targeting the three sensitive attributes, making it a more direct approach for substantially reducing their FairGap values without relying on the presence of latent sensitive information in the features. On the contrary, the Gradient Reversal method adopts a softer strategy by aiming to remove sensitive information from the learned representation, which tends to preserve predictive performance better while yielding broader fairness improvements, including for potential sensitive attributes not explicitly considered in the model, such as founder age. Overall, these two methods offer complementary strengths, Regularization is more effective when targeting fairness on specific known attributes, while Gradient Reversal provides greater generalizability and flexibility across diverse fairness concerns.

## 5.2 Future Research Directions



- **Incorporating additional sensitive attributes**

Future research could consider a broader range of potentially discriminatory factors. For example, prior literature has identified founder age as another possible source of bias in startup investment decisions (Matthews et al., 2024). However, such information is unavailable in our current dataset. Expanding data sources to include more demographic or background-related variables could allow for a more comprehensive fairness analysis. This would enable fairness-aware algorithms to address a wider spectrum of potential biases in venture funding decisions.

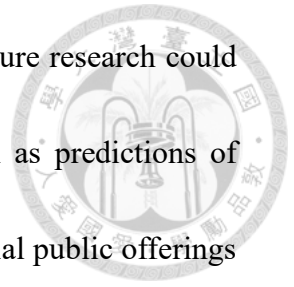
- **Algorithmic advancements to mitigate performance trade-off**

The inherent trade-off between predictive performance and fairness of the fairness-aware methods remains a key challenge, as both fairness and predictive utility are critical in decision-making scenarios. Future research could explore algorithmic innovations aimed at reducing this trade-off. This may involve reframing how fairness constraints are integrated into model objectives, or designing training frameworks that are more adaptable to varying degrees of bias during the training process, in order to better harmonize fairness and predictive performance.

- **Extending fairness research to other milestones in the startup lifecycle**

Beyond the startup early success prediction, startups may encounter different

forms and degrees of biases at various stages of their lifecycle. Future research could extend fairness-aware modeling to other critical milestones, such as predictions of follow-on funding rounds, mergers and acquisitions (M&A), or initial public offerings (IPOs). These stages often involve distinct decision-making criteria and stakeholder dynamics, which may give rise to different patterns of biases. Investigating fairness across these contexts would provide a more comprehensive understanding of how algorithmic discrimination affects startups over time, and how mitigation strategies should be adapted to each phase.



## References



Astebro, T. B. (2021). An inside peek at AI use in private equity. *Journal of Financial Data Science*, 3(3), 97-107.

Balachandra, L. (2020). How gender biases drive venture capital decision-making: exploring the gender funding gap. *Gender in Management: An International Journal*, 35(3), 261-273.

Binns, R. (2018). Fairness in machine learning: Lessons from political philosophy. *Conference on Fairness, Accountability and Transparency*, 149-159.

Bodenhausen, G. V., & Wyer, R. S. (1985). Effects of stereotypes in decision making and information-processing strategies. *Journal of Personality and Social Psychology*, 48(2), 267-282.

Bygrave, W. D., & Timmons, J. (1992). Venture capital at the crossroads. *University of Illinois at Urbana-Champaign's Academy for Entrepreneurial Leadership Historical Research Reference in Entrepreneurship*.

Caton, S., & Haas, C. (2024). Fairness in machine learning: A survey. *ACM Computing Surveys*, 56(7), 1-38.

Chen, I., Johansson, F. D., & Sontag, D. (2018). Why is my classifier discriminatory? *Advances in Neural Information Processing Systems*, 31, 3543-3554.

Chen, J., Kallus, N., Mao, X., Svacha, G., & Udell, M. (2019). Fairness under



unawareness: Assessing disparity when protected class is unobserved.

*Proceedings of the 2019 Conference on Fairness, Accountability, and*

*Transparency*, 339-348, ACM.

Corea, F., Bertinetti, G., & Cervellati, E. M. (2021). Hacking the venture industry: An

Early-stage Startups Investment framework for data-driven investors. *Machine*

*Learning with Applications*, 5, 100062.

Coval, J. D., & Moskowitz, T. J. (1999). Home bias at home: Local equity preference

in domestic portfolios. *The Journal of Finance*, 54(6), 2045-2073.

Dale, S. (2015). Heuristics and biases: The science of decision-making. *Business*

*Information Review*, 32(2), 93-99.

Dwork, C., Hardt, M., Pitassi, T., Reingold, O., & Zemel, R. (2012). Fairness through

awareness. *Proceedings of the 3rd Innovations in Theoretical Computer*

*Science Conference*, 214-226, ACM.

Ewens, M., & Townsend, R. R. (2020). Are early stage investors biased against

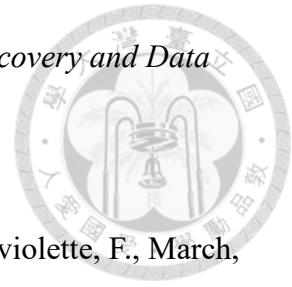
women? *Journal of Financial Economics*, 135(3), 653-677.

Fairlie, R., Robb, A., & Robinson, D. T. (2022). Black and white: Access to capital

among minority-owned start-ups. *Management Science*, 68(4), 2377-2400.

Feldman, M., Friedler, S. A., Moeller, J., Scheidegger, C., & Venkatasubramanian, S.

(2015). Certifying and removing disparate impact. *Proceedings of the 21st*



Ganin, Y., Ustinova, E., Ajakan, H., Germain, P., Larochelle, H., Laviolette, F., March, M., & Lempitsky, V. (2016). Domain-adversarial training of neural networks.

*Journal of Machine Learning Research*, 17(59), 1-35.

Gazanchyan, N. S., Hashimzade, N., Rodionova, Y., & Vershinina, N. (2017). Gender, access to finance, occupational choice, and business performance. *CESifo*

*Working Paper Series*, 6353.

Gompers, P. A., & Lerner, J. (2004). *The Venture Capital Cycle*. MIT press.

Grari, V., Ruf, B., Lamprier, S., & Detyniecki, M. (2019). Fairness-aware neural rényi minimization for continuous features. arXiv preprint arXiv:1911.04929.

Hajian, S., & Domingo-Ferrer, J. (2012). A methodology for direct and indirect discrimination prevention in data mining. *IEEE Transactions on Knowledge and Data Engineering*, 25(7), 1445-1459.

Hardt, M., Price, E., & Srebro, N. (2016). Equality of opportunity in supervised learning. *Advances in Neural Information Processing Systems*, 29.

Hellmann, T., & Puri, M. (2000). The interaction between product market and financing strategy: The role of venture capital. *The Review of Financial Studies*, 13(4), 959-984.

Hernandez, M., Raveendhran, R., Weingarten, E., & Barnett, M. (2019). How algorithms can diversify the startup pool. *MIT Sloan Management Review*, 61(1), 71-78.



Ivanitzki, T., & Rashida, J. (2023). Expand underrepresented participation in high-tech start-ups. *Proceeding of 2023 Conference for Industry and Education Collaboration*, South Carolina, ASEE.

Kamiran, F., & Calders, T. (2012). Data preprocessing techniques for classification without discrimination. *Knowledge and Information Systems*, 33(1), 1-33.

Kamishima, T., Akaho, S., Asoh, H., & Sakuma, J. (2012). Fairness-aware classifier with prejudice remover regularizer. *Proceedings of Machine Learning and Knowledge Discovery in Databases*, 7524, 35-50, Springer.

Krishna, A., Agrawal, A., & Choudhary, A. (2016). Predicting the outcome of startups: less failure, more success. *Proceedings of 2016 IEEE 16th International Conference on Data Mining Workshops*, 798-805, IEEE.

Kumar, I. E., Hines, K. E., & Dickerson, J. P. (2022). Equalizing credit opportunity in algorithms: Aligning algorithmic fairness research with us fair lending regulation. *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society*, 357-368, ACM.

Lohia, P. K., Ramamurthy, K. N., Bhide, M., Saha, D., Varshney, K. R., & Puri, R.



(2019). Bias mitigation post-processing for individual and group fairness.

*Proceedings of 2019 IEEE International Conference on Acoustics, Speech and Signal Processing*, 2847-2851, IEEE.

Matthews, M. J., Anglin, A. H., Drover, W., & Wolfe, M. T. (2024). Just a number?

Using artificial intelligence to explore perceived founder age in entrepreneurial fundraising. *Journal of Business Venturing*, 39(1), 106361.

Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, 54(6), Article 115.

Moriarty, R., Ly, H., Lan, E., & McIntosh, S. K. (2019). Deal or no deal: Predicting mergers and acquisitions at scale. *Proceedings of 2019 IEEE International Conference on Big Data*, 5552-5558, IEEE.

Paglia, J. K., & Harjoto, M. A. (2014). The effects of private equity and venture capital on sales and employment growth in small and medium-sized businesses. *Journal of Banking & Finance*, 47, 177-197.

Pedreshi, D., Ruggieri, S., & Turini, F. (2008). Discrimination-aware data mining. *Proceedings of the 14th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 560-568, ACM.

Pessach, D., & Shmueli, E. (2022). A review on fairness in machine learning. *ACM*

*Computing Surveys*, 55(3), 1-44.

QS Top Universities. (2024). *QS World University Rankings 2025*.

<https://www.topuniversities.com/world-university-rankings/2025>



Reynolds, P. D. (2016). Start-up actions and outcomes: What entrepreneurs do to reach profitability. *Foundations and Trends in Entrepreneurship*, 12(6), 443-559.

Sharchilev, B., Roizner, M., Rumyantsev, A., Ozornin, D., Serdyukov, P., & de Rijke, M. (2018). Web-based startup success prediction. *Proceedings of the 27th ACM International Conference on Information and Knowledge Management*, 2283-2291, ACM.

Te, Y.-F., Wieland, M., Frey, M., & Grabner, H. (2023a). Mitigating discriminatory biases in success prediction models for venture capitals. *Proceedings of 2023 10th IEEE Swiss Conference on Data Science*, 26-33, IEEE.

Te, Y.-F., Wieland, M., Frey, M., Pyatigorskaya, A., Schiffer, P., & Grabner, H. (2023b). Making it into a successful series a funding: An analysis of crunchbase and linkedin data. *The Journal of Finance and Data Science*, 9, 100099.

Wei, C.-P., Fang, E. S.-H., Yang, C.-S., & Liu, P.-J. (2025). To shine or not to shine: Startup success prediction by exploiting technological and venture-capital-

related features. *Information & Management*, 62(6), 104152.

Zafar, M. B., Valera, I., Ródriguez, M. G., & Gummadi, K. P. (2017). Fairness

constraints: Mechanisms for fair classification. *Artificial Intelligence and*

*Statistics*, 54, 962-970.

Zemel, R., Wu, Y., Swersky, K., Pitassi, T., & Dwork, C. (2013). Learning fair

representations. *Proceedings of International Conference on Machine*

*Learning*, 28, 325-333.



## Appendix



Table A1 presents a comprehensive list of the features employed in our methods, summarizing variables under different categories and provides their corresponding descriptions, where sensitive features are highlighted using bold formatting. It also specifies whether each feature is classified as static or dynamic, meaning it remains constant regardless of the trigger timing, or dynamic, meaning it is time-dependent and the values may vary based on the trigger event.

**Table A1: Overview of Predictive Features Used in Our Methods**

| Category | Features           | Description                                                                                                                                                                                           | Temporal Scope |
|----------|--------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------|
| Company  | <b>Region</b>      | The continent-level grouping where the focal startup is located.                                                                                                                                      | Static         |
|          | State              | The U.S. state in which the focal startup is based.                                                                                                                                                   | Static         |
|          | Industry tag       | A set of binary indicators representing the industry categories assigned to the startup. Each indicator corresponds to a specific tag, such as Advertising, Health Care, Information Technology, etc. | Static         |
|          | Industry tag group | Industry tags are grouped into six clusters based on their co-occurrence patterns. For each startup, the total number of tags it holds in each group is computed.                                     | Static         |
|          | Age                | Number of months from the startup's founding date to the trigger round.                                                                                                                               | Dynamic        |

|          |                                    |                                                                                                                                                                                                           |         |
|----------|------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|
| Founders | Number of advisors                 | Number of advisors involved prior to the trigger event.                                                                                                                                                   | Dynamic |
|          | Number of founders                 | Total number of individuals in the founding team.                                                                                                                                                         | Static  |
|          | <b>Female founder ratio</b>        | Proportion of female individuals among all founders.                                                                                                                                                      | Static  |
|          | <b>Non-Caucasian founder ratio</b> | Proportion of non-Caucasian individuals among all founders.                                                                                                                                               | Static  |
|          | Subject average                    | Average number of distinct subjects studied across all founders.                                                                                                                                          | Static  |
|          | Subject                            | A set of values representing the average number of founders who have studied in specific academic subject areas. Each value corresponds to a subject tag, such as Business, Engineering, Humanities, etc. | Static  |
|          | Degree average                     | Average number of academic degrees earned by all founders in the startup team.                                                                                                                            | Static  |
|          | Degree                             | A set of values representing the average number of degrees earned by founder team members at each education level (e.g., Bachelor, Master, PhD, or High School and below).                                | Static  |
|          | Top 100 university average         | Average number of founding team members who obtained degrees from the top 100 universities based on the latest QS World University Rankings.                                                              | Static  |
|          | Job average                        | Average number of previous jobs held by each founding team member.                                                                                                                                        | Static  |
|          | Job title                          | A set of values representing the average number of previous jobs held by founding team members under each job title category (e.g.,                                                                       | Static  |

|                            |                                    |                                                                                                                                                                                   |         |
|----------------------------|------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|
|                            |                                    | CEO, CTO, Engineering & Tech, etc.).                                                                                                                                              |         |
|                            | Job type                           | A set of values representing the average number of previous jobs held by founding team members under each job type category (e.g., Advisor, Board Member, Executive, etc.).       | Static  |
|                            | Work years average                 | Average years of work experience among founding team members.                                                                                                                     | Static  |
|                            | Serial entrepreneur average        | Average number of previous companies founded by team members.                                                                                                                     | Static  |
| Investment-Funding Details | Trigger type                       | Indicates whether the trigger round is an angel or seed investment.                                                                                                               | Dynamic |
|                            | Trigger round index                | The sequence number of the trigger round within the company's overall funding history.                                                                                            | Dynamic |
|                            | Number of funding rounds           | Total number of funding rounds the startup received.                                                                                                                              | Dynamic |
|                            | Total funding amount               | Total amount of capital raised (in USD) by the startup.                                                                                                                           | Dynamic |
|                            | Number of funding rounds (by type) | A set of values representing the number of funding rounds the startup has received for each funding type considered possible before Series A (e.g., angel, pre-seed, seed, etc.). | Dynamic |
|                            | Total funding amount (by type)     | A set of values representing the total amount of capital raised (in USD) by the startup for each funding type considered possible before Series A.                                | Dynamic |
|                            | Burn rate                          | Startup's capital consumption rate, calculated by dividing total capital raised by company age (in months).                                                                       | Dynamic |

|                                 |                                                   |                                                                                                                                                                     |         |
|---------------------------------|---------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|
|                                 | Capital concentration rate                        | Degree of capital concentration across all previous funding rounds, reflecting the extent to which a small number of investors contributed the majority of funding. | Dynamic |
| Investment-Historical Investors | Number of investors                               | Total number of historical investors who have participated in funding rounds up to the trigger round.                                                               | Dynamic |
|                                 | Average investors                                 | Average number of investors in the past funding rounds of the focal startup.                                                                                        | Dynamic |
|                                 | Number of distinct investors                      | Total number of distinct investors who have participated in funding rounds up to the trigger round.                                                                 | Dynamic |
|                                 | Average investments                               | Average number of past investments made by the historical investors who have invested in the focal startup.                                                         | Dynamic |
|                                 | Average investments amount                        | Average total amount (in USD) previously invested by the historical investors.                                                                                      | Dynamic |
|                                 | Average number of distinct organizations invested | Average number of distinct organizations previously invested in by the historical investors.                                                                        | Dynamic |
|                                 | Number of institutional investors                 | Number of institutional (organization-type) investors among all prior investors.                                                                                    | Dynamic |
|                                 | Number of individual investors                    | Number of individual (person-type) investors among all prior investors.                                                                                             | Dynamic |
|                                 | Number of distinct institutional investors        | Number of distinct institutional investors among all prior investors.                                                                                               | Dynamic |
|                                 | Number of distinct individual investors           | Number of distinct individual investors among all prior investors.                                                                                                  | Dynamic |

|                            |                                                                                                                                                        |         |
|----------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------|---------|
| Average success rate       | Average success rate of the historical investors.                                                                                                      | Dynamic |
| Max success rate           | Maximum success rate of the historical investors.                                                                                                      | Dynamic |
| Total network centrality   | Total network centrality of all historical investors, separately measured by betweenness centrality (BC), degree centrality (DC), and PageRank (PR).   | Dynamic |
| Average network centrality | Average network centrality of all historical investors, separately measured by betweenness centrality (BC), degree centrality (DC), and PageRank (PR). | Dynamic |

Table A2 provides the descriptive statistics of our dataset, classified into two categories by the label of target variable, failure and success, indicating whether a startup successfully secured Series A funding or not during the observation period.

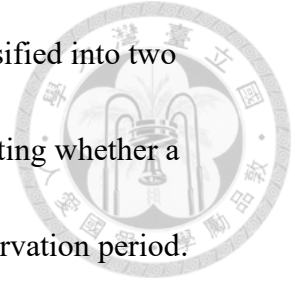


Table A2: Descriptive Statistics of Our Dataset by Startup Success Label

| Feature                                    | Stats (label = failure)                                                                                                                      | Stats (label = success)                                                                                                                    |
|--------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------|
| Region                                     | North America: 4,879 (54.37%); Europe: 2,310 (25.74%); East Asia: 203 (2.26%); Other: 1,582 (17.63%);                                        | North America: 2,908 (62.97%); Europe: 925 (20.03%); East Asia: 140 (3.03%); Other: 645 (13.97%)                                           |
| State                                      | non-US: 4,432 (49.39%); CA: 2,159 (24.06%); NY: 637 (7.10%); MA: 205 (2.28%); TX: 161 (1.79%); FL: 128 (1.43%); other states: 1,252 (13.95%) | non-US: 1,860 (40.28%); CA: 1,342 (29.06%); NY: 487 (10.55%); MA: 146 (3.16%); TX: 105 (2.27%); FL: 49 (1.06%); other states: 629 (13.62%) |
| Trigger type                               | seed: 8,172 (91.06%); angel: 802 (8.94%)                                                                                                     | seed: 4,423 (95.78%); angel: 195 (4.22%)                                                                                                   |
| Industry tag_Administrative Services       | False: 8,435 (93.99%); True: 539 (6.01%)                                                                                                     | False: 4,319 (93.53%); True: 299 (6.47%)                                                                                                   |
| Industry tag_Advertising                   | False: 8,544 (95.21%); True: 430 (4.79%)                                                                                                     | False: 4,502 (97.49%); True: 116 (2.51%)                                                                                                   |
| Industry tag_Agriculture and Farming       | False: 8,867 (98.81%); True: 107 (1.19%)                                                                                                     | False: 4,552 (98.57%); True: 66 (1.43%)                                                                                                    |
| Industry tag_Apps                          | False: 7,135 (79.51%); True: 1,839 (20.49%)                                                                                                  | False: 3,857 (83.52%); True: 761 (16.48%)                                                                                                  |
| Industry tag_Artificial Intelligence (AI)  | False: 6,955 (77.50%); True: 2,019 (22.50%)                                                                                                  | False: 3,269 (70.79%); True: 1,349 (29.21%)                                                                                                |
| Industry tag_Biotechnology                 | False: 8,834 (98.44%); True: 140 (1.56%)                                                                                                     | False: 4,505 (97.55%); True: 113 (2.45%)                                                                                                   |
| Industry tag_Blockchain and Cryptocurrency | False: 8,593 (95.75%); True: 381 (4.25%)                                                                                                     | False: 4,433 (95.99%); True: 185 (4.01%)                                                                                                   |
| Industry tag_Clothing and Apparel          | False: 8,873 (98.87%); True: 101 (1.13%)                                                                                                     | False: 4,581 (99.20%); True: 37 (0.80%)                                                                                                    |
| Industry tag_Commerce and Shopping         | False: 7,839 (87.35%); True: 1,135 (12.65%)                                                                                                  | False: 4,087 (88.50%); True: 531 (11.50%)                                                                                                  |
| Industry tag_Community and Lifestyle       | False: 8,572 (95.52%); True: 402 (4.48%)                                                                                                     | False: 4,458 (96.54%); True: 160 (3.46%)                                                                                                   |
| Industry tag_Consumer Electronics          | False: 8,349 (93.04%); True: 625 (6.96%)                                                                                                     | False: 4,318 (93.50%); True: 300 (6.50%)                                                                                                   |
| Industry tag_Consumer Goods                | False: 8,800 (98.06%); True: 174 (1.94%)                                                                                                     | False: 4,525 (97.99%); True: 93 (2.01%)                                                                                                    |

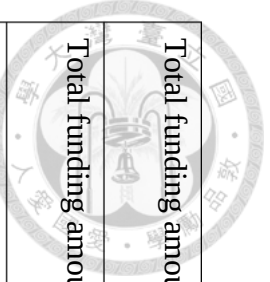
|                                               |                                             |                                             |
|-----------------------------------------------|---------------------------------------------|---------------------------------------------|
| Industry tag_Content and Publishing           | False: 8,606 (95.90%); True: 368 (4.10%)    | False: 4,490 (97.23%); True: 128 (2.77%)    |
| Industry tag_Data and Analytics               | False: 6,035 (67.25%); True: 2,939 (32.75%) | False: 2,753 (59.61%); True: 1,865 (40.39%) |
| Industry tag_Design                           | False: 8,500 (94.72%); True: 474 (5.28%)    | False: 4,387 (95.00%); True: 231 (5.00%)    |
| Industry tag_Education                        | False: 8,228 (91.69%); True: 746 (8.31%)    | False: 4,282 (92.72%); True: 336 (7.28%)    |
| Industry tag_Energy                           | False: 8,866 (98.80%); True: 108 (1.20%)    | False: 4,544 (98.40%); True: 74 (1.60%)     |
| Industry tag_Events                           | False: 8,827 (98.36%); True: 147 (1.64%)    | False: 4,562 (98.79%); True: 56 (1.21%)     |
| Industry tag_Financial Services               | False: 7,701 (85.81%); True: 1,273 (14.19%) | False: 3,745 (81.10%); True: 873 (18.90%)   |
| Industry tag_Food and Beverage                | False: 8,754 (97.55%); True: 220 (2.45%)    | False: 4,476 (96.93%); True: 142 (3.07%)    |
| Industry tag_Gaming                           | False: 8,721 (97.18%); True: 253 (2.82%)    | False: 4,531 (98.12%); True: 87 (1.88%)     |
| Industry tag_Government and Military          | False: 8,882 (98.97%); True: 92 (1.03%)     | False: 4,553 (98.59%); True: 65 (1.41%)     |
| Industry tag_Hardware                         | False: 7,038 (78.43%); True: 1,936 (21.57%) | False: 3,723 (80.62%); True: 895 (19.38%)   |
| Industry tag_Health Care                      | False: 8,156 (90.88%); True: 818 (9.12%)    | False: 4,087 (88.50%); True: 531 (11.50%)   |
| Industry tag_Information Technology           | False: 6,356 (70.83%); True: 2,618 (29.17%) | False: 3,183 (68.93%); True: 1,435 (31.07%) |
| Industry tag_Internet Services                | False: 6,464 (72.03%); True: 2,510 (27.97%) | False: 3,658 (79.21%); True: 960 (20.79%)   |
| Industry tag_Lending and Investments          | False: 8,742 (97.41%); True: 232 (2.59%)    | False: 4,433 (95.99%); True: 185 (4.01%)    |
| Industry tag_Manufacturing                    | False: 8,753 (97.54%); True: 221 (2.46%)    | False: 4,455 (96.47%); True: 163 (3.53%)    |
| Industry tag_Media and Entertainment          | False: 7,654 (85.29%); True: 1,320 (14.71%) | False: 4,257 (92.18%); True: 361 (7.82%)    |
| Industry tag_Messaging and Telecommunications | False: 8,637 (96.24%); True: 337 (3.76%)    | False: 4,475 (96.90%); True: 143 (3.10%)    |
| Industry tag_Mobile                           | False: 7,191 (80.13%); True: 1,783 (19.87%) | False: 3,821 (82.74%); True: 797 (17.26%)   |
| Industry tag_Music and Audio                  | False: 8,810 (98.17%); True: 164 (1.83%)    | False: 4,573 (99.03%); True: 45 (0.97%)     |
| Industry tag_Natural Resources                | False: 8,904 (99.22%); True: 70 (0.78%)     | False: 4,582 (99.22%); True: 36 (0.78%)     |

|                                           |                                                      |                                                      |
|-------------------------------------------|------------------------------------------------------|------------------------------------------------------|
| Industry tag_Navigation and Mapping       | False: 8,809 (98.16%); True: 165 (1.84%)             | False: 4,553 (98.59%); True: 65 (1.41%)              |
| Industry tag_Other                        | False: 7,477 (83.32%); True: 1,497 (16.68%)          | False: 3,846 (83.28%); True: 772 (16.72%)            |
| Industry tag_Payments                     | False: 8,397 (93.57%); True: 577 (6.43%)             | False: 4,234 (91.68%); True: 384 (8.32%)             |
| Industry tag_Platforms                    | False: 8,658 (96.48%); True: 316 (3.52%)             | False: 4,522 (97.92%); True: 96 (2.08%)              |
| Industry tag_Privacy and Security         | False: 8,516 (94.90%); True: 458 (5.10%)             | False: 4,280 (92.68%); True: 338 (7.32%)             |
| Industry tag_Professional Services        | False: 8,159 (90.92%); True: 815 (9.08%)             | False: 4,158 (90.04%); True: 460 (9.96%)             |
| Industry tag_Real Estate                  | False: 8,558 (95.36%); True: 416 (4.64%)             | False: 4,395 (95.17%); True: 223 (4.83%)             |
| Industry tag_Sales and Marketing          | False: 7,849 (87.46%); True: 1,125 (12.54%)          | False: 4,139 (89.63%); True: 479 (10.37%)            |
| Industry tag_Science and Engineering      | False: 6,855 (76.39%); True: 2,119 (23.61%)          | False: 3,226 (69.86%); True: 1,392 (30.14%)          |
| Industry tag_Social Impact                | False: 8,908 (99.26%); True: 66 (0.74%)              | False: 4,595 (99.50%); True: 23 (0.50%)              |
| Industry tag_Sports                       | False: 8,723 (97.20%); True: 251 (2.80%)             | False: 4,526 (98.01%); True: 92 (1.99%)              |
| Industry tag_Sustainability               | False: 8,868 (98.82%); True: 106 (1.18%)             | False: 4,524 (97.96%); True: 94 (2.04%)              |
| Industry tag_Transportation               | False: 8,369 (93.26%); True: 605 (6.74%)             | False: 4,208 (91.12%); True: 410 (8.88%)             |
| Industry tag_Travel and Tourism           | False: 8,713 (97.09%); True: 261 (2.91%)             | False: 4,502 (97.49%); True: 116 (2.51%)             |
| Industry tag_Video                        | False: 8,659 (96.49%); True: 315 (3.51%)             | False: 4,508 (97.62%); True: 110 (2.38%)             |
| Industry tag_group_Finance                | $\mu = 0.75$ , $\sigma = 1.01$ , max = 5, min = 0    | $\mu = 0.8$ , $\sigma = 1.04$ , max = 5, min = 0     |
| Industry tag_group_Lifestyle              | $\mu = 0.26$ , $\sigma = 0.5$ , max = 3, min = 0     | $\mu = 0.31$ , $\sigma = 0.54$ , max = 3, min = 0    |
| Industry tag_group_MediaTech              | $\mu = 1.29$ , $\sigma = 1.44$ , max = 8, min = 0    | $\mu = 0.95$ , $\sigma = 1.25$ , max = 8, min = 0    |
| Industry tag_group_Hardware               | $\mu = 0.36$ , $\sigma = 0.65$ , max = 3, min = 0    | $\mu = 0.32$ , $\sigma = 0.62$ , max = 3, min = 0    |
| Industry tag_group_Production / GreenTech | $\mu = 0.32$ , $\sigma = 0.6$ , max = 5, min = 0     | $\mu = 0.42$ , $\sigma = 0.68$ , max = 5, min = 0    |
| Industry tag_group_High-tech              | $\mu = 0.99$ , $\sigma = 1.09$ , max = 5, min = 0    | $\mu = 1.19$ , $\sigma = 1.16$ , max = 5, min = 0    |
| Age                                       | $\mu = 19.13$ , $\sigma = 12.64$ , max = 48, min = 0 | $\mu = 18.75$ , $\sigma = 11.94$ , max = 48, min = 0 |

|                             |                                                              |                                                              |
|-----------------------------|--------------------------------------------------------------|--------------------------------------------------------------|
| Number of advisors          | $\mu = 0.08, \sigma = 0.31, \text{max} = 4, \text{min} = 0$  | $\mu = 0.09, \sigma = 0.33, \text{max} = 3, \text{min} = 0$  |
| Number of founders          | $\mu = 1.78, \sigma = 0.86, \text{max} = 8, \text{min} = 1$  | $\mu = 2.07, \sigma = 0.96, \text{max} = 10, \text{min} = 1$ |
| Female founder ratio        | $\mu = 0.11, \sigma = 0.27, \text{max} = 1, \text{min} = 0$  | $\mu = 0.09, \sigma = 0.23, \text{max} = 1, \text{min} = 0$  |
| Non-Caucasian founder ratio | $\mu = 0.3, \sigma = 0.4, \text{max} = 1, \text{min} = 0$    | $\mu = 0.3, \sigma = 0.39, \text{max} = 1, \text{min} = 0$   |
| Subject average             | $\mu = 0.49, \sigma = 0.47, \text{max} = 3, \text{min} = 0$  | $\mu = 0.56, \sigma = 0.45, \text{max} = 2, \text{min} = 0$  |
| Subject_Biology/Health      | $\mu = 0.01, \sigma = 0.1, \text{max} = 2, \text{min} = 0$   | $\mu = 0.01, \sigma = 0.08, \text{max} = 1, \text{min} = 0$  |
| Subject_Business            | $\mu = 0.14, \sigma = 0.31, \text{max} = 3, \text{min} = 0$  | $\mu = 0.15, \sigma = 0.31, \text{max} = 2, \text{min} = 0$  |
| Subject_CS/IT               | $\mu = 0.15, \sigma = 0.32, \text{max} = 3, \text{min} = 0$  | $\mu = 0.17, \sigma = 0.32, \text{max} = 2, \text{min} = 0$  |
| Subject_Engineering         | $\mu = 0.07, \sigma = 0.23, \text{max} = 3, \text{min} = 0$  | $\mu = 0.09, \sigma = 0.24, \text{max} = 2, \text{min} = 0$  |
| Subject_Humanities          | $\mu = 0.02, \sigma = 0.12, \text{max} = 2, \text{min} = 0$  | $\mu = 0.01, \sigma = 0.1, \text{max} = 1, \text{min} = 0$   |
| Subject_Law                 | $\mu = 0.01, \sigma = 0.08, \text{max} = 2, \text{min} = 0$  | $\mu = 0.01, \sigma = 0.07, \text{max} = 1, \text{min} = 0$  |
| Subject_Math/Physics        | $\mu = 0.02, \sigma = 0.11, \text{max} = 2, \text{min} = 0$  | $\mu = 0.02, \sigma = 0.12, \text{max} = 1, \text{min} = 0$  |
| Subject_Media/Comm          | $\mu = 0.01, \sigma = 0.07, \text{max} = 1, \text{min} = 0$  | $\mu = 0.01, \sigma = 0.07, \text{max} = 2, \text{min} = 0$  |
| Subject_Other               | $\mu = 0.05, \sigma = 0.19, \text{max} = 2, \text{min} = 0$  | $\mu = 0.06, \sigma = 0.2, \text{max} = 2, \text{min} = 0$   |
| Subject_SocialSci           | $\mu = 0.04, \sigma = 0.16, \text{max} = 2, \text{min} = 0$  | $\mu = 0.04, \sigma = 0.17, \text{max} = 2, \text{min} = 0$  |
| Degree average              | $\mu = 0.51, \sigma = 0.5, \text{max} = 4, \text{min} = 0$   | $\mu = 0.58, \sigma = 0.48, \text{max} = 3, \text{min} = 0$  |
| Degree_Bachelor             | $\mu = 0.24, \sigma = 0.4, \text{max} = 3, \text{min} = 0$   | $\mu = 0.27, \sigma = 0.39, \text{max} = 3, \text{min} = 0$  |
| Degree_High school or below | $\mu = 0.04, \sigma = 0.18, \text{max} = 3, \text{min} = 0$  | $\mu = 0.03, \sigma = 0.15, \text{max} = 2, \text{min} = 0$  |
| Degree_Master               | $\mu = 0.19, \sigma = 0.39, \text{max} = 4, \text{min} = 0$  | $\mu = 0.22, \sigma = 0.38, \text{max} = 3, \text{min} = 0$  |
| Degree_PhD                  | $\mu = 0.04, \sigma = 0.18, \text{max} = 2, \text{min} = 0$  | $\mu = 0.06, \sigma = 0.2, \text{max} = 2, \text{min} = 0$   |
| Top 100 university average  | $\mu = 0.1, \sigma = 0.28, \text{max} = 2, \text{min} = 0$   | $\mu = 0.15, \sigma = 0.31, \text{max} = 2, \text{min} = 0$  |
| Job average                 | $\mu = 1.87, \sigma = 1.78, \text{max} = 22, \text{min} = 0$ | $\mu = 2.08, \sigma = 1.97, \text{max} = 31, \text{min} = 0$ |

|                                      |                                                              |                                                               |
|--------------------------------------|--------------------------------------------------------------|---------------------------------------------------------------|
| Job title_CEO                        | $\mu = 0.54, \sigma = 0.62, \text{max} = 6, \text{min} = 0$  | $\mu = 0.5, \sigma = 0.59, \text{max} = 4, \text{min} = 0$    |
| Job title_CFO                        | $\mu = 0.01, \sigma = 0.07, \text{max} = 1, \text{min} = 0$  | $\mu = 0.01, \sigma = 0.08, \text{max} = 3, \text{min} = 0$   |
| Job title_CIO                        | $\mu = 0.01, \sigma = 0.07, \text{max} = 2, \text{min} = 0$  | $\mu = 0.0045, \sigma = 0.06, \text{max} = 1, \text{min} = 0$ |
| Job title_CMO                        | $\mu = 0.01, \sigma = 0.11, \text{max} = 3, \text{min} = 0$  | $\mu = 0.01, \sigma = 0.11, \text{max} = 3, \text{min} = 0$   |
| Job title_COO                        | $\mu = 0.06, \sigma = 0.21, \text{max} = 3, \text{min} = 0$  | $\mu = 0.05, \sigma = 0.18, \text{max} = 2, \text{min} = 0$   |
| Job title_CPO                        | $\mu = 0.01, \sigma = 0.1, \text{max} = 2, \text{min} = 0$   | $\mu = 0.02, \sigma = 0.09, \text{max} = 1, \text{min} = 0$   |
| Job title_CSO                        | $\mu = 0.01, \sigma = 0.06, \text{max} = 1, \text{min} = 0$  | $\mu = 0.01, \sigma = 0.07, \text{max} = 1, \text{min} = 0$   |
| Job title_CTO                        | $\mu = 0.31, \sigma = 0.54, \text{max} = 5, \text{min} = 0$  | $\mu = 0.34, \sigma = 0.55, \text{max} = 6, \text{min} = 0$   |
| Job title_Education / Research       | $\mu = 0.03, \sigma = 0.19, \text{max} = 4, \text{min} = 0$  | $\mu = 0.05, \sigma = 0.21, \text{max} = 3, \text{min} = 0$   |
| Job title_Engineering / Tech         | $\mu = 0.13, \sigma = 0.39, \text{max} = 6, \text{min} = 0$  | $\mu = 0.15, \sigma = 0.4, \text{max} = 6, \text{min} = 0$    |
| Job title_Executive / Management     | $\mu = 0.22, \sigma = 0.65, \text{max} = 19, \text{min} = 0$ | $\mu = 0.31, \sigma = 0.71, \text{max} = 10, \text{min} = 0$  |
| Job title_External Advisor           | $\mu = 0.18, \sigma = 0.57, \text{max} = 12, \text{min} = 0$ | $\mu = 0.21, \sigma = 0.61, \text{max} = 21, \text{min} = 0$  |
| Job title_Finance / HR / Admin       | $\mu = 0.01, \sigma = 0.1, \text{max} = 3, \text{min} = 0$   | $\mu = 0.01, \sigma = 0.08, \text{max} = 2, \text{min} = 0$   |
| Job title_Marketing / Sales          | $\mu = 0.04, \sigma = 0.2, \text{max} = 4, \text{min} = 0$   | $\mu = 0.04, \sigma = 0.21, \text{max} = 4, \text{min} = 0$   |
| Job title_Operations / PM / Customer | $\mu = 0.08, \sigma = 0.29, \text{max} = 4, \text{min} = 0$  | $\mu = 0.1, \sigma = 0.3, \text{max} = 3, \text{min} = 0$     |
| Job title_Other                      | $\mu = 0.16, \sigma = 0.44, \text{max} = 14, \text{min} = 0$ | $\mu = 0.2, \sigma = 0.51, \text{max} = 8, \text{min} = 0$    |
| Job title_Product / Design / Content | $\mu = 0.07, \sigma = 0.28, \text{max} = 5, \text{min} = 0$  | $\mu = 0.08, \sigma = 0.29, \text{max} = 4, \text{min} = 0$   |
| Job type_advisor                     | $\mu = 0.11, \sigma = 0.46, \text{max} = 9, \text{min} = 0$  | $\mu = 0.14, \sigma = 0.51, \text{max} = 12, \text{min} = 0$  |
| Job type_Board member                | $\mu = 0.12, \sigma = 0.51, \text{max} = 17, \text{min} = 0$ | $\mu = 0.17, \sigma = 0.56, \text{max} = 12, \text{min} = 0$  |
| Job type_Board observer              | $\mu = 0, \sigma = 0.06, \text{max} = 4, \text{min} = 0$     | $\mu = 0.01, \sigma = 0.08, \text{max} = 2, \text{min} = 0$   |
| Job type_Employee                    | $\mu = 0.45, \sigma = 0.8, \text{max} = 8, \text{min} = 0$   | $\mu = 0.55, \sigma = 0.93, \text{max} = 10, \text{min} = 0$  |
| Job type_Executive                   | $\mu = 1.19, \sigma = 1.03, \text{max} = 16, \text{min} = 0$ | $\mu = 1.21, \sigma = 1.01, \text{max} = 11, \text{min} = 0$  |

|                                                |                                                                          |                                                                          |
|------------------------------------------------|--------------------------------------------------------------------------|--------------------------------------------------------------------------|
| Work years average                             | $\mu = 3.86, \sigma = 4.76, \max = 57, \min = 0$                         | $\mu = 4.2, \sigma = 4.89, \max = 58, \min = 0$                          |
| Serial entrepreneur average                    | $\mu = 0.47, \sigma = 0.78, \max = 10, \min = 0$                         | $\mu = 0.46, \sigma = 0.71, \max = 8, \min = 0$                          |
| Trigger round index                            | $\mu = 1.43, \sigma = 0.76, \max = 7, \min = 1$                          | $\mu = 1.46, \sigma = 0.75, \max = 7, \min = 1$                          |
| Number of funding rounds                       | $\mu = 1.92, \sigma = 1.21, \max = 12, \min = 1$                         | $\mu = 1.94, \sigma = 1.16, \max = 9, \min = 1$                          |
| Total funding amount                           | $\mu = 1,108,486.57, \sigma = 1,891,753.22, \max = 40,828,073, \min = 0$ | $\mu = 2,278,130.04, \sigma = 2,672,658.78, \max = 58,710,343, \min = 0$ |
| Number of funding rounds_Angel                 | $\mu = 0.17, \sigma = 0.46, \max = 5, \min = 0$                          | $\mu = 0.11, \sigma = 0.39, \max = 5, \min = 0$                          |
| Number of funding rounds_Convertible note      | $\mu = 0.05, \sigma = 0.25, \max = 4, \min = 0$                          | $\mu = 0.05, \sigma = 0.25, \max = 3, \min = 0$                          |
| Number of funding rounds_Corporate round       | $\mu = 0.0009, \sigma = 0.03, \max = 1, \min = 0$                        | $\mu = 0.0009, \sigma = 0.03, \max = 1, \min = 0$                        |
| Number of funding rounds_Debt financing        | $\mu = 0.02, \sigma = 0.14, \max = 3, \min = 0$                          | $\mu = 0.02, \sigma = 0.17, \max = 4, \min = 0$                          |
| Number of funding rounds_Equity crowdfunding   | $\mu = 0.01, \sigma = 0.1, \max = 6, \min = 0$                           | $\mu = 0.0017, \sigma = 0.04, \max = 1, \min = 0$                        |
| Number of funding rounds_Grant                 | $\mu = 0.06, \sigma = 0.31, \max = 6, \min = 0$                          | $\mu = 0.06, \sigma = 0.31, \max = 4, \min = 0$                          |
| Number of funding rounds_Non equity assistance | $\mu = 0.05, \sigma = 0.25, \max = 3, \min = 0$                          | $\mu = 0.05, \sigma = 0.23, \max = 4, \min = 0$                          |
| Number of funding rounds_Pre seed              | $\mu = 0.19, \sigma = 0.52, \max = 5, \min = 0$                          | $\mu = 0.2, \sigma = 0.46, \max = 4, \min = 0$                           |
| Number of funding rounds_Product crowdfunding  | $\mu = 0.0047, \sigma = 0.07, \max = 2, \min = 0$                        | $\mu = 0.01, \sigma = 0.07, \max = 1, \min = 0$                          |
| Number of funding rounds_Seed                  | $\mu = 1.37, \sigma = 0.85, \max = 7, \min = 0$                          | $\mu = 1.44, \sigma = 0.79, \max = 6, \min = 0$                          |
| Total funding amount_Angel                     | $\mu = 59,916.05, \sigma = 356,634.76, \max = 16,884,057, \min = 0$      | $\mu = 62,246.42, \sigma = 345,121.51, \max = 7,100,000, \min = 0$       |



|                                            |                                                                                 |                                                                                   |
|--------------------------------------------|---------------------------------------------------------------------------------|-----------------------------------------------------------------------------------|
| Total funding amount_Convertible note      | $\mu = 11,498.29$ , $\sigma = 142,796.55$ , $\max = 5,000,000$ , $\min = 0$     | $\mu = 14,487.33$ , $\sigma = 157,810.11$ , $\max = 4,478,000$ , $\min = 0$       |
| Total funding amount_Corporate round       | $\mu = 3,413.72$ , $\sigma = 301,121.56$ , $\max = 28,505,763$ , $\min = 0$     | $\mu = 0$ , $\sigma = 0$ , $\max = 0$ , $\min = 0$                                |
| Total funding amount_Debt financing        | $\mu = 8,479.93$ , $\sigma = 198,251.17$ , $\max = 15,250,000$ , $\min = 0$     | $\mu = 26,762.67$ , $\sigma = 633,935.39$ , $\max = 33,157,402$ , $\min = 0$      |
| Total funding amount_Equity crowdfunding   | $\mu = 3,294.99$ , $\sigma = 96,817.42$ , $\max = 7,966,138$ , $\min = 0$       | $\mu = 953.57$ , $\sigma = 29,591.17$ , $\max = 1,215,465$ , $\min = 0$           |
| Total funding amount_Grant                 | $\mu = 11,065.38$ , $\sigma = 140,183.99$ , $\max = 8,000,000$ , $\min = 0$     | $\mu = 16,585.55$ , $\sigma = 152,259.19$ , $\max = 3,793,696$ , $\min = 0$       |
| Total funding amount_Non equity assistance | $\mu = 2,202.88$ , $\sigma = 97,000.56$ , $\max = 6,353,366$ , $\min = 0$       | $\mu = 95.89$ , $\sigma = 2,879.46$ , $\max = 157,547$ , $\min = 0$               |
| Total funding amount_Pre seed              | $\mu = 42,674.81$ , $\sigma = 207,855.62$ , $\max = 4,000,000$ , $\min = 0$     | $\mu = 65,215.06$ , $\sigma = 317,217.66$ , $\max = 5,790,935$ , $\min = 0$       |
| Total funding amount_Product crowdfunding  | $\mu = 1,966.33$ , $\sigma = 45,632.15$ , $\max = 1,700,000$ , $\min = 0$       | $\mu = 2,482.43$ , $\sigma = 57,296.48$ , $\max = 2,000,000$ , $\min = 0$         |
| Total funding amount_Seed                  | $\mu = 963,974.21$ , $\sigma = 1,736,210.61$ , $\max = 40,828,073$ , $\min = 0$ | $\mu = 2,089,301.11$ , $\sigma = 2,505,161.72$ , $\max = 58,710,343$ , $\min = 0$ |
| Burn rate                                  | $\mu = 137,260.79$ , $\sigma = 1,094,544.98$ , $\max = 54,720,000$ , $\min = 0$ | $\mu = 240,098.94$ , $\sigma = 1,226,937.83$ , $\max = 48,640,000$ , $\min = 0$   |
| Capital concentration rate                 | $\mu = 0.42$ , $\sigma = 0.39$ , $\max = 1$ , $\min = 0$                        | $\mu = 0.34$ , $\sigma = 0.33$ , $\max = 1$ , $\min = 0$                          |
| Number of investors                        | $\mu = 3.61$ , $\sigma = 3.7$ , $\max = 40$ , $\min = 0$                        | $\mu = 5.48$ , $\sigma = 5.29$ , $\max = 84$ , $\min = 0$                         |

|                                                   |                                                                                            |                                                                                              |
|---------------------------------------------------|--------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------|
| Average investors                                 | $\mu = 2.2, \sigma = 2.02, \text{max} = 32, \text{min} = 0$                                | $\mu = 3.24, \sigma = 3, \text{max} = 84, \text{min} = 0$                                    |
| Number of distinct investors                      | $\mu = 2.83, \sigma = 3.02, \text{max} = 41, \text{min} = 0$                               | $\mu = 4.42, \sigma = 4.22, \text{max} = 55, \text{min} = 0$                                 |
| Average investments                               | $\mu = 89.75, \sigma = 201.02, \text{max} = 2,281, \text{min} = 0$                         | $\mu = 84.55, \sigma = 155.52, \text{max} = 2,036, \text{min} = 0$                           |
| Average investments amount                        | $\mu = 52,574,882.46, \sigma = 170,560,764.46, \text{max} = 5,470,119,343, \text{min} = 0$ | $\mu = 106,094,246.35, \sigma = 328,435,604.97, \text{max} = 10,734,190,065, \text{min} = 0$ |
| Average number of distinct organizations invested | $\mu = 71.4, \sigma = 156.8, \text{max} = 1,677, \text{min} = 0$                           | $\mu = 65.31, \sigma = 120.87, \text{max} = 1,504, \text{min} = 0$                           |
| Number of institutional investors                 | $\mu = 2.64, \sigma = 2.72, \text{max} = 26, \text{min} = 0$                               | $\mu = 3.98, \sigma = 3.64, \text{max} = 40, \text{min} = 0$                                 |
| Number of individual investors                    | $\mu = 0.89, \sigma = 2.03, \text{max} = 28, \text{min} = 0$                               | $\mu = 1.5, \sigma = 2.97, \text{max} = 64, \text{min} = 0$                                  |
| Number of distinct institutional investors        | $\mu = 2.36, \sigma = 2.3, \text{max} = 23, \text{min} = 0$                                | $\mu = 3.54, \sigma = 2.99, \text{max} = 28, \text{min} = 0$                                 |
| Number of distinct individual investors           | $\mu = 0.84, \sigma = 1.89, \text{max} = 24, \text{min} = 0$                               | $\mu = 1.43, \sigma = 2.79, \text{max} = 64, \text{min} = 0$                                 |
| Average success rate                              | $\mu = 0.07, \sigma = 0.1, \text{max} = 1, \text{min} = 0$                                 | $\mu = 0.1, \sigma = 0.11, \text{max} = 1, \text{min} = 0$                                   |
| Max success rate                                  | $\mu = 0.15, \sigma = 0.21, \text{max} = 1, \text{min} = 0$                                | $\mu = 0.25, \sigma = 0.25, \text{max} = 1, \text{min} = 0$                                  |
| Total network centrality_BC                       | $\mu = 0.0046, \sigma = 0.0103, \text{max} = 0.1068, \text{min} = 0$                       | $\mu = 0.0072, \sigma = 0.0131, \text{max} = 0.1326, \text{min} = 0$                         |
| Total network centrality_DC                       | $\mu = 0.1953, \sigma = 0.4229, \text{max} = 5.1404, \text{min} = 0$                       | $\mu = 0.3529, \sigma = 0.6061, \text{max} = 5.6413, \text{min} = 0$                         |
| Total network centrality_PR                       | $\mu = 0.001, \sigma = 0.0019, \text{max} = 0.026, \text{min} = 0$                         | $\mu = 0.0017, \sigma = 0.0027, \text{max} = 0.0324, \text{min} = 0$                         |
| Average network centrality_BC                     | $\mu = 0.0013, \sigma = 0.0037, \text{max} = 0.0397, \text{min} = 0$                       | $\mu = 0.0013, \sigma = 0.0028, \text{max} = 0.0397, \text{min} = 0$                         |
| Average network centrality_DC                     | $\mu = 0.0508, \sigma = 0.1201, \text{max} = 1, \text{min} = 0$                            | $\mu = 0.0611, \sigma = 0.1093, \text{max} = 1, \text{min} = 0$                              |



| Average network centrality_PR | $\mu = 0.0003, \sigma = 0.0005, \max = 0.0049, \min = 0$ | $\mu = 0.0003, \sigma = 0.0005, \max = 0.0053, \min = 0$ |
|-------------------------------|----------------------------------------------------------|----------------------------------------------------------|
|-------------------------------|----------------------------------------------------------|----------------------------------------------------------|