

# **TÀI LIỆU HƯỚNG DẪN THỰC HIỆN DỰ ÁN NGHIÊN CỨU KHOA HỌC**

**Người thực hiện : ThS. Đặng Thiện Tâm**

## **MỤC LỤC**

|                                                                                                            |           |
|------------------------------------------------------------------------------------------------------------|-----------|
| <b>PHẦN I. GIỚI THIỆU CHUNG VỀ MÔ HÌNH HỒI QUY TUYẾN TÍNH BỘI.....</b>                                     | <b>2</b>  |
| 1.1 Giới thiệu mô hình hồi quy tuyến tính bội .....                                                        | 2         |
| 1.1.1 Mô hình hồi quy đa tuyến tuyến tính có dạng phương trình .....                                       | 2         |
| 1.1.2 Ví dụ minh họa: Mô hình kinh tế lượng mô tả các yếu tố quyết định tiền lương của người lao động..... | 3         |
| 1.2 Các giai đoạn thiết lập nghiên cứu phân tích mô hình hồi quy tuyến tính bội 4                          |           |
| 1.2.1 Xác định thang đo.....                                                                               | 5         |
| 1.2.2 Các xác định kích thước mẫu trong nghiên cứu .....                                                   | 8         |
| 1.2.3 Xây dựng bảng hỏi và thu thập dữ liệu .....                                                          | 13        |
| 1.2.4 Phân tích độ tin cậy Cronbach's Alpha .....                                                          | 13        |
| 1.2.5 Phân tích nhân tố khám phá EFA (Exploratory Factor Analysis).....                                    | 16        |
| 1.2.6 Phân tích hồi quy đa biến .....                                                                      | 19        |
| <b>PHẦN II. DỰ ÁN NGHIÊN CỨU : CẤU TRÚC ĐỀ CƯƠNG CHUNG CỦA MỘT DỰ ÁN NGHIÊN CỨU .....</b>                  | <b>27</b> |
| Chương 1: Tổng quan về nghiên cứu (10 giờ làm việc) .....                                                  | 27        |
| Chương 2: Cơ sở lý thuyết và mô hình nghiên cứu (10 giờ làm việc) .....                                    | 27        |
| Chương 3. Phương pháp nghiên cứu (10 giờ làm việc).....                                                    | 29        |
| Chương 4 Kết quả nghiên cứu (10 giờ làm việc).....                                                         | 30        |
| Chương 5: Kết luận và kiến nghị (5 giờ làm việc) .....                                                     | 31        |
| <b>PHẦN III: HƯỚNG DẪN ĐỌC Ý NGHĨA VÀ CÁC BƯỚC THỰC HIỆN NCKH .....</b>                                    | <b>33</b> |
| 1. Phương pháp hồi quy đa biến sử dụng trên SPSS.....                                                      | 33        |
| 2. Hướng dẫn phương pháp nghiên cứu định lượng cho dự án khoa học đối với ngôn ngữ R và Smart PLS 4 .....  | 51        |

## **PHẦN I. GIỚI THIỆU CHUNG VỀ MÔ HÌNH HỒI QUY TUYẾN TÍNH BỘI**

### **1.1 Giới thiệu mô hình hồi quy tuyến tính bội**

Các phương pháp kinh tế lượng có liên quan trong hầu hết mọi ngành kinh tế học ứng dụng. Chúng phát huy tác dụng hoặc khi chúng ta có một lý thuyết kinh tế để kiểm tra hoặc khi chúng ta nghĩ đến một mối quan hệ có tầm quan trọng nhất định đối với các quyết định kinh doanh hoặc phân tích chính sách. Một phân tích thực nghiệm sử dụng dữ liệu để kiểm tra một lý thuyết hoặc để ước tính một mối quan hệ. Một trong những mô hình được sử dụng nhất trong NCKH đó là mô hình phân tích hồi quy đa tuyến tính. Hồi quy đa tuyến tính (linear regression) là một phương pháp thống kê để dự đoán giá trị của một biến phụ thuộc (hay còn gọi là biến mục tiêu) dựa trên một hoặc nhiều biến độc lập (hay còn gọi là biến đầu vào). Trong hồi quy đa tuyến tính, mối quan hệ giữa biến mục tiêu và biến đầu vào được mô hình hóa thành một phương trình tuyến tính. Hồi quy đa tuyến tính có thể được áp dụng trong nhiều lĩnh vực, bao gồm kinh tế học, khoa học dữ liệu, thống kê, và nhiều lĩnh vực khác. Việc áp dụng hồi quy đa tuyến tính giúp giải quyết nhiều vấn đề trong thực tế, như dự báo giá cổ phiếu, dự đoán giá nhà đất, hoặc phân tích mối quan hệ giữa các yếu tố trong kinh tế học.

#### **1.1.1 Mô hình hồi quy đa tuyến tính có dạng phương trình**

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_k X_k + u$$

Trong đó:

◆ Y: biến phụ thuộc

- ♦  $X_1, X_2, X_k$ : biến độc lập
- ♦  $u$ : Độ lệch chuẩn - Biểu thị sự khác biệt giữa giá trị dự đoán và giá trị thực tế của biến mục tiêu.
- ♦  $\beta_0, \beta_1, \beta_k$ : là các hệ số hồi quy (hay còn gọi là trọng số), biểu thị mức độ ảnh hưởng của mỗi biến đầu vào đến biến mục tiêu

**Phương trình ước lượng:**  $\hat{Y} = \hat{\beta}_0 + \hat{\beta}_1 X_1 + \hat{\beta}_2 X_2 + \hat{u}$

- $\hat{Y}$  biến dự đoán,  $\hat{\beta}_0, \hat{\beta}_1$  là hệ số dự đoán
- $\hat{u} = y - \hat{y}$
- Hệ số  $\hat{\beta}_1, \hat{\beta}_2$ , đo lường mức độ thay đổi của biến phụ thuộc khi biến độc lập thay đổi một đơn vị

**Giải thích các hệ số :**  $\hat{\beta}_k = \frac{\Delta y}{\Delta x_k}$

- Hệ số  $\hat{\beta}_k$  đo lường mức độ thay đổi của biến phụ thuộc khi biến độc lập  $x$  tăng một đơn vị trong khi giữ cố định các nhân tố khác.
- Giả định rằng các nhân tố không quan sát được không thay đổi nếu biến độc lập thay đổi,  $\frac{\Delta u}{\Delta x_k} = 0$

### 1.1.2 Ví dụ minh họa: Mô hình kinh tế lượng mô tả các yếu tố quyết định tiền lương của người lao động.

Mô hình hồi quy tuyến tính bội:

|                                                                                                            |
|------------------------------------------------------------------------------------------------------------|
| $\text{Wage} = \beta_0 + \beta_1 \text{education} + \beta_2 \text{experience} + \beta_3 \text{tenure} + u$ |
|------------------------------------------------------------------------------------------------------------|

- ♦ Biến phụ thuộc: Tiền lương (wage)
- ♦  $X_1, X_2, X_3$ : biến độc lập : Giáo dục (education), Kinh nghiệm(experience) và thâm niên (tenure)
- ♦  $\beta_0, \beta_1, \beta_k$ : là các hệ số được ước tính, biểu thị mức độ ảnh hưởng của mỗi biến độc lập ( Giáo dục, thâm niên kinh nghiệm đến biến phụ thuộc ( tiền lương).

Giá trị các biến: Tiền lương được tính bằng đô la/giờ, trình độ học vấn tính bằng năm, kinh nghiệm tính bằng năm, năm làm việc trong công ty (thâm niên)

**Phương trình ước tính cho giá trị dự đoán của giá trị:**

$$\widehat{Wage} = \beta_0 + \beta_1 \text{education} + \beta_2 \text{experience} + \beta_3 \text{tenure}$$

- $\beta_1$  đo lường sự thay đổi của tiền lương liên quan đến việc học thêm một năm khi các yếu tố khác cố định
- $\beta_2$  đo lường sự thay đổi về tiền lương khi có thêm một năm kinh nghiệm khi các yếu tố khác không đổi
- $\beta_3$  đo lường sự thay đổi của tiền lương liên quan đến một năm nhiệm kỳ khi giữ nguyên các yếu tố khác

**Kết quả phương trình ước lượng**

$$\widehat{Wage} = -2.78 + 0.6 \text{ education} + 0.03 \text{ experience} + 0.54 \text{ tenure}$$

**Kết quả dự đoán**

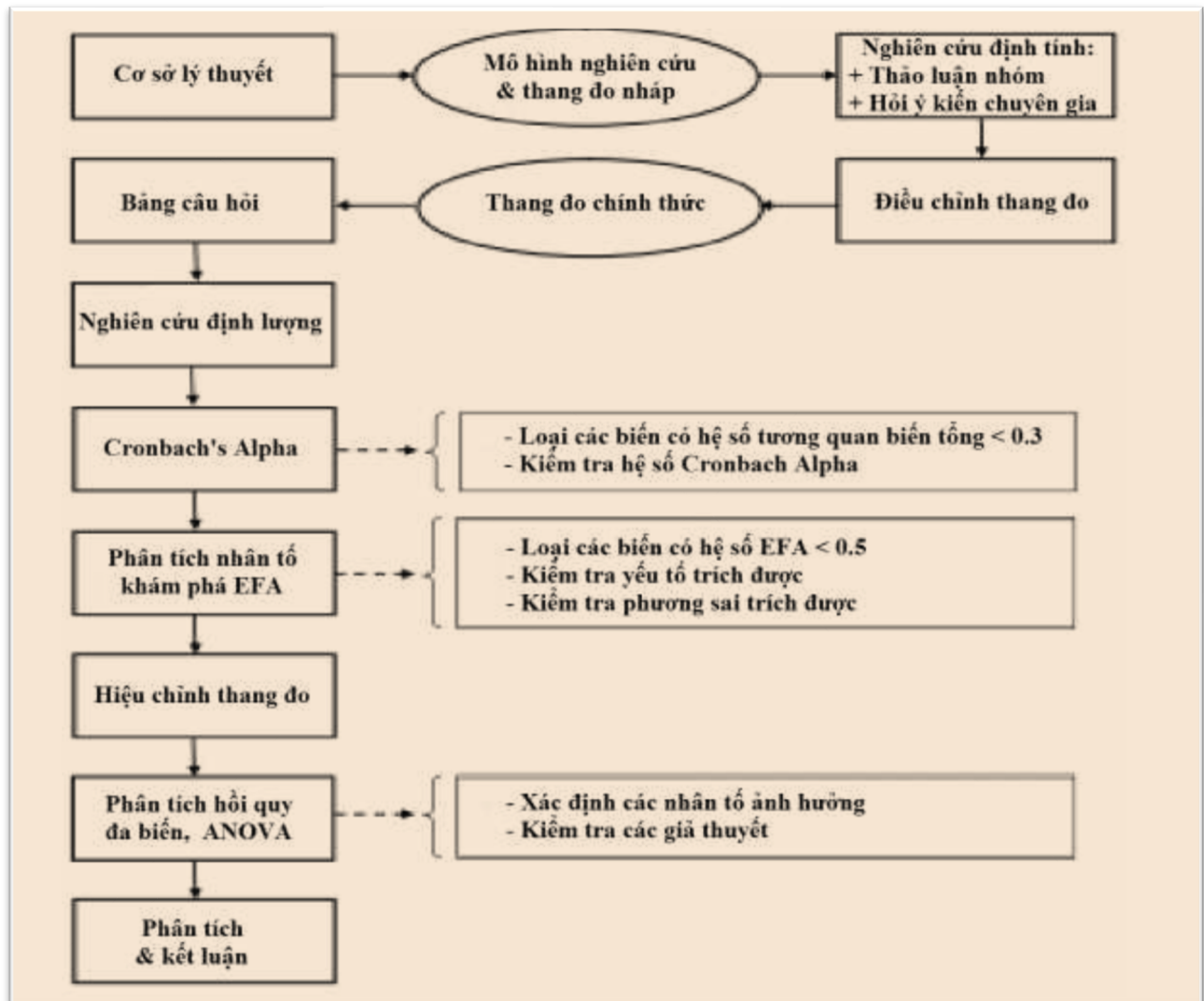
- $\beta_1$ : Tiền lương theo giờ tăng thêm 0.6 đô la cho mỗi năm đào tạo bổ sung, giữ nguyên các yếu tố khác.
- $\beta_2$ : Tiền lương theo giờ tăng thêm 0,03 đô la cho mỗi năm kinh nghiệm bổ sung, giữ nguyên các yếu tố khác.
- $\beta_3$ : Tiền lương theo giờ tăng thêm 0,45 đô la cho mỗi năm làm việc bổ sung trong công ty, giữ nguyên các yếu tố khác.
- Nếu tất cả các biến giải thích bằng 0, thì tiền lương của một người là - \$2.78

## **1.2 Các giai đoạn thiết lập nghiên cứu phân tích mô hình hồi quy tuyến tính bội**

Thiết lập nghiên cứu phân tích mô hình hồi quy tuyến tính bội là quá trình quan trọng để nghiên cứu mối quan hệ giữa các biến trong một hệ thống. Thiết lập nghiên cứu phân tích mô hình hồi quy tuyến tính bội đòi hỏi các bước cẩn thận và kỹ lưỡng

để đảm bảo tính chính xác và độ tin cậy của kết quả. Để thiết lập một nghiên cứu phân tích mô hình hồi quy tuyến tính bội, cần thực hiện các bước sau đây:

Hình 1.2 Các giai đoạn thực hiện một nghiên cứu định lượng phân tích mô hình hồi quy tuyến tính bội



### 1.2.1 Xác định thang đo

Trong nghiên cứu khoa học, các thang đo được sử dụng để đo lường các biến số liên quan đến đối tượng nghiên cứu. Việc xác định thang đo phù hợp là rất quan trọng để đảm bảo tính chính xác và đáng tin cậy của kết quả nghiên cứu.

Có ba loại thang đo phổ biến được sử dụng trong nghiên cứu khoa học:

- (1) Thang đo định tính (nominal scale): Thang đo định tính được sử dụng để đo lường các biến số có tính chất phân loại, không có giá trị số chính xác, ví dụ như giới tính, chủng tộc, tình trạng hôn nhân, ...
- (2) Thang đo định lượng: Thang đo định lượng được sử dụng để đo lường các biến số có giá trị số chính xác. Có hai loại thang đo định lượng chính là:
  - + Thang đo định lượng rời rạc: Sử dụng để đo lường các biến số có giá trị số rời rạc (không liên tục) như số lượng, số lần, số người, số lượng sản phẩm,...
  - + Thang đo định lượng liên tục: Sử dụng để đo lường các biến số có giá trị số liên tục như chiều cao, cân nặng, nhiệt độ, thời gian,...
- (3) Thang đo tỷ lệ: Thang đo tỷ lệ được sử dụng để đo lường các biến số có tính chất tỷ lệ, có giá trị số chính xác và có một giá trị tối thiểu. Ví dụ như số lượng tiền, số lượng sản phẩm bán ra,...
- (4) Khi lựa chọn thang đo, nghiên cứu viên cần phải xác định đúng loại thang đo phù hợp với biến số nghiên cứu, đảm bảo tính chính xác và độ tin cậy của dữ liệu thu được.

Trong khuôn khổ nghiên cứu khảo sát thang điểm Li-kert từ 1 điểm (thể hiện ý kiến cho rằng họ có mức kỳ vọng không nhiều hoặc mức hài lòng rất thấp) đến 5 điểm (thể hiện mức kỳ vọng rất cao hoặc mức độ rất hài lòng) là thang đo được sử dụng phổ biến.

Ngoài ra chúng ta có thể sử dụng Likert 7 điểm, trong đó 1: Hoàn toàn phản đối và 7 hoàn toàn đồng ý. (Strongly Disagree, Disagree, Somewhat Disagree, Somewhat agree, Agree, Strongly agree).

**Theo Bollen (2011), có 3 loại thang đo đo lường gồm:**

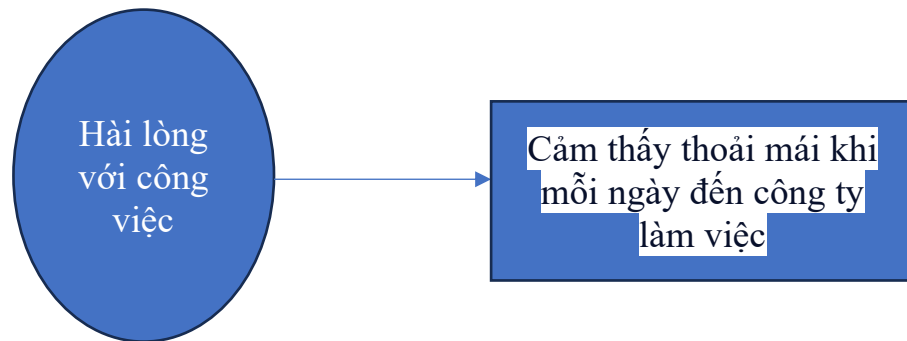
**+ Effect Indicator (Reflective Measurement): thang đo kết quả (Reflective)**

**Nội dung cần nắm về Reflective:**

1. Biến quan sát là kết quả được tạo ra từ biến tiềm ẩn

2. Các biến quan sát có tương quan cùng chiều với nhau ở mức khá, mạnh
3. Khi bỏ 1 biến quan sát khỏi mô hình thang đo, biến tiềm ẩn không bị ảnh hưởng nhiều
4. Các biến quan sát có thể thay thế vai trò cho nhau
5. Có thể sử dụng các phân tích: Cronbach Alpha, Composite Reliability, EFA, CFA

Ví dụ HL1: Cảm thấy thoải mái khi mỗi ngày đến công ty làm việc

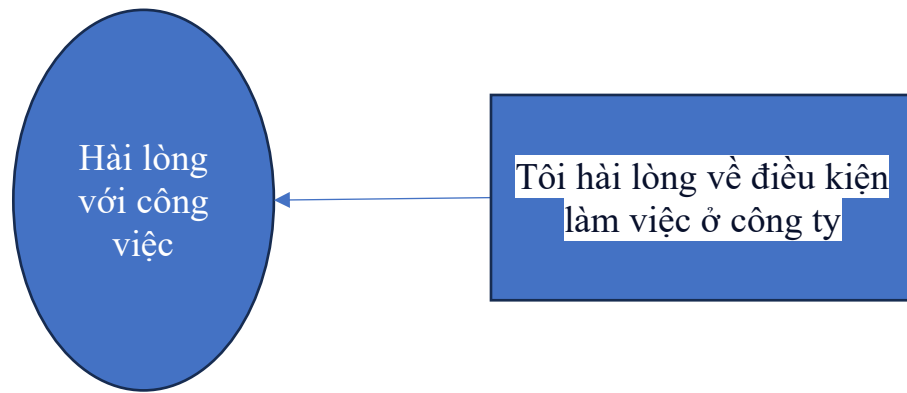


#### **+ Composite Indicator (Formative Measurement): Thang đo nguyên nhân**

##### **Nội dung cần nắm về Formative:**

1. Biến quan sát là các thành phần cấu thành nên biến tiềm ẩn
2. Các biến quan sát ít có sự tương quan với nhau
3. Khi bỏ 1 biến quan sát khỏi mô hình thang đo, biến tiềm ẩn bị ảnh hưởng rất nhiều
4. Các biến quan sát không thể thay thế vai trò cho nhau
5. **KHÔNG THỂ** sử dụng các phân tích: Cronbach Alpha, Composite Reliability, EFA, CFA

Ví dụ HL1: Tôi hài lòng về điều kiện làm việc ở công ty



### + Causal Indicator (Causal Measurement): Thang đo Causal

#### 1.2.2 Các xác định kích thước mẫu trong nghiên cứu

##### 1.2.2.1 Yếu tố ảnh hưởng tới quyết định chọn cỡ mẫu

Kích thước mẫu (cỡ mẫu) của nghiên cứu càng lớn, sai số trong các ước lượng sẽ càng thấp, khả năng đại diện cho tổng thể càng cao. Tuy nhiên, việc thu thập cỡ mẫu lớn sẽ làm tiêu tốn nhiều thời gian, công sức, tiền bạc ở toàn bộ các khâu từ thu thập, kiểm tra, phân tích. Do đó việc chọn kích thước mẫu cần phải được xem xét một cách có cân nhắc để mọi thứ được cân bằng và hiệu quả. Sự lựa chọn cỡ mẫu sẽ phụ thuộc vào:

- Độ tin cậy cần có của dữ liệu. Nghĩa là mức độ chắc chắn rằng các đặc điểm của cỡ mẫu được chọn phải khái quát được cho đặc điểm tổng thể.
- Sai số mà nghiên cứu có thể chấp nhận được. Đó là độ chính xác chúng ta yêu cầu cho bất kỳ ước lượng được thực hiện trên mẫu.
- Các loại kiểm định, phân tích sẽ thực hiện. Một số kỹ thuật thống kê yêu cầu cỡ mẫu phải đạt một ngưỡng nhất định thì các ước lượng mới có ý nghĩa.
- Kích thước của tổng thể. Mẫu nghiên cứu sẽ cần chiếm một tỷ lệ nhất định so với kích thước của tổng thể.

### 1.2.2.2. Xác định cỡ mẫu theo ước lượng tổng thể

Theo Yamane Taro (1967), việc xác định kích thước mẫu sẽ được chia làm hai trường hợp: không biết tổng thể và biết được tổng thể.

#### *a. Trường hợp không biết quy mô tổng thể*

Chúng ta sẽ sử dụng công thức sau:

$$n = Z^2 \times \frac{p \times (1 - p)}{e^2}$$

Trong đó:

- $n$ : kích thước mẫu cần xác định.
- $Z$ : giá trị tra bảng phân phối Z dựa vào độ tin cậy lựa chọn. Thông thường, độ tin cậy được sử dụng là 95% tương ứng với  $Z = 1.96$ .
- $p$ : tỷ lệ ước lượng cỡ mẫu n thành công. Thường chúng ta chọn  $p = 0.5$  để tích số  $p(1-p)$  là lớn nhất, điều này đảm bảo an toàn cho mẫu n ước lượng.
- $e$ : sai số cho phép. Thường ba tỷ lệ sai số hay sử dụng là:  $\pm 0.01$  (1%),  $\pm 0.05$  (5%),  $\pm 0.1$  (10%), trong đó mức phổ biến nhất là  $\pm 0.05$ .

Ví dụ: Nghiên cứu sự hài lòng của khách hàng đã dùng sản phẩm nước giải khát Pepsi Cola tại TP.HCM. Đây là tổng thể không xác định được quy mô vì chúng ta không biết được có bao nhiêu khách hàng đã uống nước Pepsi Cola ở TP.HCM. Như vậy cỡ mẫu tối thiểu cần có của nghiên cứu sẽ là 385 người:

$$n = 1.96^2 \times \frac{0.5 \times (1 - 0.5)}{0.05^2} = 384.16$$

***b. Trường hợp biết quy mô tổng thể***

Chúng ta sẽ sử dụng công thức sau:

$$n = \frac{N}{1 + N \times e^2}$$

Trong đó:

- $n$ : kích thước mẫu cần xác định.
- $N$ : quy mô tổng thể.
- $e$ : sai số cho phép. Thường ba tỷ lệ sai số hay sử dụng là:  $\pm 0.01$  (1%),  $\pm 0.05$  (5%),  $\pm 0.1$  (10%), trong đó mức phổ biến nhất là  $\pm 0.05$ .

Ví dụ: Nghiên cứu sự hài lòng của khách hàng đã mua sữa bột Ensure Gold trong tháng 8 năm 2020 tại siêu thị Coopmart Phú Thọ (Quận 11, TP.HCM). Siêu thị tổng hợp danh sách khách hàng từ hệ thống thì có 1000 khách hàng, đây là tổng thể xác định được quy mô. Như vậy cỡ mẫu tối thiểu cần có của nghiên cứu nếu sai số  $e = \pm 0.05$  sẽ là 286 người:

$$n = \frac{1000}{1 + 1000 \times 0.05^2} = 285.71$$

### **1.2.2.2. Xác định cỡ mẫu theo ước lượng tổng thể**

Việc xác định cỡ mẫu theo ước lượng tổng thể thường yêu cầu cỡ mẫu lớn. Tuy nhiên, nhà nghiên cứu lại có quỹ thời gian giới hạn và nếu không có nguồn tài chính tài trợ thì khả năng lấy mẫu theo ước lượng tổng thể sẽ khó có thể thực hiện. Do đó, các nhà nghiên cứu thường sử dụng công thức lấy mẫu dựa vào phương pháp định lượng được sử dụng để phân tích dữ liệu. Hai phương pháp yêu cầu cỡ mẫu lớn thường là hồi quy và phân tích nhân tố khám phá (EFA).

#### ***a. Kích thước mẫu theo EFA***

Theo Hair và cộng sự (2014)[1], kích thước mẫu tối thiểu để sử dụng EFA là 50, tốt hơn là từ 100 trở lên. Tỷ lệ số quan sát trên một biến phân tích là 5:1 hoặc 10:1, một số nhà nghiên cứu cho rằng tỷ lệ này nên là 20:1. “Số quan sát” hiểu một cách đơn giản là số phiếu khảo sát hợp lệ cần thiết; “biến đo lường” là một câu hỏi đo lường trong bảng khảo sát. Ví dụ, nếu bảng khảo sát của chúng ta có 30 câu hỏi sử dụng thang đo Likert 5 mức độ (tương ứng với 30 biến quan sát thuộc các nhân tố khác nhau), 30 câu này được sử dụng để phân tích trong một lần EFA. Áp dụng tỷ lệ 5:1, cỡ mẫu tối thiểu sẽ là  $30 \times 5 = 150$ , nếu tỷ lệ 10:1 thì cỡ mẫu tối thiểu là  $30 \times 10 = 300$ . Kích thước mẫu này lớn hơn kích thước tối thiểu 50 hoặc 100, vì vậy chúng ta cần cỡ mẫu tối thiểu để thực hiện phân tích nhân tố khám phá EFA là 150 hoặc 300 tùy tỷ lệ lựa chọn dựa trên khả năng có thể khảo sát được.

#### ***b. Kích thước mẫu theo hồi quy***

Đối với kích thước mẫu tối thiểu cho phân tích hồi quy, Green (1991)[2] đưa ra hai trường hợp. Trường hợp một, nếu mục đích phép hồi quy chỉ đánh giá mức độ phù hợp tổng quát của mô hình như  $R^2$ , kiểm định F ... thì cỡ mẫu tối thiểu là  $50 + 8m$  (m là số lượng biến độc lập hay còn gọi là predictor tham gia vào hồi quy). Trường hợp hai, nếu mục đích muốn đánh giá các yếu tố của từng biến độc lập như

kiểm định t, hệ số hồi quy ... thì cỡ mẫu tối thiểu nên là  $104 + m$  ( $m$  là số lượng biến độc lập). Lưu ý rằng,  $m$  là số biến độc lập chúng ta đưa vào phân tích hồi quy, không phải là số biến quan sát hay số câu hỏi của nghiên cứu. Giả sử chúng ta xây dựng bảng khảo sát gồm 4 biến độc lập (4 thang đo), mỗi thang đo biến độc lập này được đo lường bằng 5 câu hỏi Likert (5 biến quan sát), như vậy tổng cộng chúng ta có 20 biến quan sát. Sau bước phân tích EFA, 4 thang đo này vẫn giữ nguyên như lý thuyết ban đầu, điều này đồng nghĩa có 4 biến độc lập sẽ được sử dụng cho phân tích hồi quy, tức  $m = 4$  không phải  $m = 20$ .

Harris (1985)[3] cho rằng cỡ mẫu phù hợp để chạy hồi quy đa biến phải bằng số biến độc lập cộng thêm ít nhất là 50. Ví dụ, phép hồi quy có 4 biến độc lập tham gia, thì cỡ mẫu tối thiểu phải là  $4 + 50 = 54$ . Hair và cộng sự (2014)[4] cho rằng cỡ mẫu tối thiểu nên theo tỷ lệ 5:1, tức là 5 quan sát cho một biến độc lập. Như vậy, nếu có 4 biến độc lập tham gia vào hồi quy, cỡ mẫu tối thiểu sẽ là  $5 \times 4 = 20$ . Tuy nhiên, 5:1 chỉ là cỡ mẫu tối thiểu cần đạt, để kết quả hồi quy có ý nghĩa thống kê cao hơn, cỡ mẫu lý tưởng nên theo tỷ lệ 10:1 hoặc 15:1. Riêng với trường hợp sử dụng phương pháp đưa biến vào lần lượt Stepwise trong hồi quy, cỡ mẫu nên theo tỷ lệ 50:1.

Nếu một bài nghiên cứu sử dụng kết hợp nhiều phương pháp xử lý thì sẽ lấy kích thước mẫu cần thiết lớn nhất trong các phương pháp. Ví dụ, nếu bài nghiên cứu vừa sử dụng phân tích EFA và vừa phân tích hồi quy. Kích thước mẫu cần thiết của EFA là 200, kích thước mẫu cần thiết của hồi quy là 100, chúng ta sẽ chọn kích thước mẫu cần thiết của nghiên cứu là 200 hoặc từ 200 trở lên. Thường chúng ta sử dụng phân tích EFA cùng với phân tích hồi quy trong cùng một bài luận văn, một bài nghiên cứu. EFA luôn đòi hỏi cỡ mẫu lớn hơn rất nhiều so với hồi quy, chính vì vậy chúng ta có thể sử dụng công thức tính kích thước mẫu tối thiểu cho EFA làm công thức tính kích thước mẫu cho nghiên cứu. Cũng lưu ý rằng, đây là cỡ mẫu tối thiểu,

nếu chúng ta sử dụng cỡ mẫu lớn hơn kích thước tối thiểu, nghiên cứu sẽ càng có giá trị.

### **1.2.3 Xây dựng bảng hỏi và thu thập dữ liệu**

Bảng hỏi là một công cụ phổ biến được sử dụng trong việc thu thập dữ liệu từ người tham gia nghiên cứu. Để xây dựng bảng hỏi và thu thập dữ liệu, cần thực hiện các bước sau:

- (1) Xác định mục tiêu và nội dung của nghiên cứu: Xác định mục tiêu nghiên cứu và những thông tin cần thu thập.
- (2) Xác định đối tượng nghiên cứu: Xác định đối tượng nghiên cứu, tức là những người mà bạn muốn thu thập thông tin từ họ.
- (3) Xác định kích thước mẫu: Xác định kích thước mẫu cần thiết để đảm bảo độ chính xác và độ đại diện của dữ liệu.
- (4) Thiết kế bảng hỏi: Thiết kế bảng hỏi bao gồm việc chọn loại câu hỏi và cách trình bày chúng để thu thập thông tin.
- (5) Kiểm tra tính đầy đủ và tính phù hợp: Kiểm tra tính đầy đủ và tính phù hợp của bảng hỏi bằng cách đưa ra một số câu hỏi thử và thu thập phản hồi.
- (6) Thực hiện thu thập dữ liệu: Thực hiện việc thu thập dữ liệu từ người tham gia nghiên cứu bằng cách phát hành bảng hỏi và thu thập phản hồi.
- (7) Kiểm tra và xử lý dữ liệu: Kiểm tra và xử lý dữ liệu thu thập được để đảm bảo tính chính xác và độ tin cậy của dữ liệu.

Tóm lại, việc xây dựng bảng hỏi và thu thập dữ liệu là một bước quan trọng trong quá trình nghiên cứu khoa học. Để đạt được kết quả nghiên cứu chính xác và tin cậy, cần phải thực hiện các bước này một cách cẩn thận và chuyên nghiệp.

### **1.2.4 Phân tích độ tin cậy Cronbach's Alpha**

Phân tích độ tin cậy Cronbach's Alpha là giai đoạn đầu tiên trong phân tích hồi quy tuyến tính bội sau khi chúng ta hoàn tất giai đoạn khảo sát dữ liệu trong nghiên cứu. Phân tích độ tin cậy Cronbach's alpha cũng giúp cho nhà nghiên cứu có thể xác

định được những câu hỏi hoặc biến nào trong bộ dữ liệu có độ tin cậy cao và nên được sử dụng để đo lường một đặc tính nào đó, và những câu hỏi hoặc biến nào có độ tin cậy thấp và nên được loại bỏ khỏi bộ dữ liệu để đảm bảo tính đáng tin cậy và chính xác của kết quả nghiên cứu.

Cronbach's alpha là thước đo độ tin cậy nhất quán nội bộ của một thang đo hoặc bảng câu hỏi. Nó thường được sử dụng trong tâm lý học, giáo dục và khoa học xã hội để đánh giá độ tin cậy của một tập hợp các mục nhằm đo lường cùng một cấu trúc hoặc lĩnh vực. Hệ số của Cronbach là một phép kiểm định thống kê về mức độ chặt chẽ mà các mục hỏi trong thang đo tương quan với nhau.

Cronbach's alpha được tính bằng cách so sánh điểm số của từng hạng mục thang đo với tổng điểm của mỗi lần quan sát (thường là người trả lời khảo sát hoặc người kiểm tra riêng lẻ), sau đó so sánh điểm số đó với phương sai cho tất cả điểm số của từng hạng mục:

$$\alpha = \left( \frac{k}{k-1} \right) \left( 1 - \frac{\sum_{i=1}^k \sigma_{y_i}^2}{\sigma_x^2} \right)$$

Trong đó:

- $K$ : đề cập đến số lượng các biến quan sát
- $\sigma_{y_i}^2$ : đề cập đến phương sai (variance) liên quan đến biến quan sát  $i$
- $\sigma_x^2$ : đề cập đến phương sai (variance) liên quan đến tổng số quan sát

Cronbach's alpha nằm trong khoảng từ 0 đến 1, với các giá trị càng cao cho thấy tính nhất quán bên trong càng lớn. Alpha bằng 0 cho biết không có sự thống nhất bên trong giữa các mục, trong khi alpha bằng 1 biểu thị sự nhất quán bên trong hoàn hảo.

Vì hệ số Cronbach Alpha chỉ là giới hạn dưới của độ tin cậy của thang đo (Theo GS.TS Nguyễn Đình Thọ), và còn nhiều đại lượng đo lường độ tin cậy, độ giá trị của thang đo, nên ở giai đoạn đầu khi xây dựng bảng câu hỏi, hệ số này nằm trong phạm vi từ 0.6 đến 0.8 là chấp nhận được.

Đây là phương pháp cho phép người phân tích loại bỏ các biến không phù hợp và hạn chế các biến rác trong quá trình nghiên cứu và đánh giá độ tin cậy của thang đo bằng hệ số thông qua hệ số Cronbach Alpha. Những biến có hệ số tương quan biến tổng (item-total correlation) nhỏ hơn 0.3 sẽ bị loại. Thang đo có hệ số Cronbach Alpha từ 0.6 trở lên là có thể sử dụng được trong trường hợp khái niệm đang nghiên cứu mới (Nunnally, 1978; Peterson, 1994; Slater, 1995). Thông thường, thang đo có Cronbach Alpha từ 0.7 đến 0.8 là sử dụng được. Nhiều nhà nghiên cứu cho rằng khi thang đo có độ tin cậy từ 0.8 trở lên đến gần 1 là thang đo lường tốt.

Khi xuất hiện giá trị Cronbach's Alpha âm có nghĩa là giá trị này bị âm là do xảy ra hiện tượng trung bình hiệp phương sai âm giữa các biến quan sát. Giả định độ tin cậy thang đo đang bị vi phạm. Bạn cần kiểm tra lại các biến quan sát. Hiểu một cách đơn giản là thang đo hoàn toàn không có độ tin cậy, các biến quan sát trong nhóm không có sự tương quan gì với nhau, hoặc các biến quan sát đa hướng, ngược chiều nhau. Những hướng xử lý với trường hợp Cronbach Alpha âm như sau:

- (1) Kiểm tra lại dữ liệu được nhập có chính xác chưa, có bị lỗi nhầm giá trị, có bị lỗi nhập sai cột, sai hàng không.
- (2) Khi đã kiểm tra rồi và kết quả không có sai sót, hãy thử xem lại bảng câu hỏi xem có hiện tượng ngược chiều thang đo không. Nếu có, bắt buộc phải điều chỉnh bảng câu hỏi và khảo sát lại.
- (3) Kiểm tra lại bảng câu hỏi, việc giá trị các biến quan sát thu thập được gần như không có mối tương quan với nhau dẫn đến **Cronbach's Alpha bị âm** lý do chính xuất phát từ việc thiết kế các biến quan sát bên trong nhân tố không hợp

lý. Đối chiếu với các cơ sở lý luận, các bảng câu hỏi cùng đề tài, tham khảo ý kiến giảng viên hướng dẫn để điều chỉnh lại bảng câu hỏi khảo sát cho hợp lý.

- (4) Trường hợp bảng khảo sát cũng đã tốt rồi, có thể vấn đề đến từ việc hợp tác của đáp viên hoặc do thực tế thang đo này hoàn toàn không có độ tin cậy với dữ liệu thực nghiệm, lúc này chúng ta có hai lựa chọn, hoặc là điều chỉnh lại thang đo, chọn lọc đối tượng khảo sát để tiến hành khảo sát lại hoặc sẽ kết luận thang đo không có độ tin cậy.

### **1.2.5 Phân tích nhân tố khám phá EFA (Exploratory Factor Analysis)**

Phân tích nhân tố khám phá (EFA) là một phương pháp phân tích nhân tố được sử dụng để khám phá mối quan hệ giữa các biến và xác định số lượng các nhân tố ẩn. EFA thường được sử dụng để giảm số lượng biến trong một tập dữ liệu lớn bằng cách tìm ra các nhân tố ẩn để thay thế cho những biến ban đầu. Phương pháp này rất có ích cho việc xác định các tập hợp biến cần thiết cho vấn đề nghiên cứu và được sử dụng để tìm mối quan hệ giữa các biến với nhau.

Các chỉ số cần lưu ý trong phân tích nhân tố khám phá:

- *Trị số KMO (Kaiser-Meyer - Olkin)*: KMO là một phương pháp được sử dụng để đánh giá tính phù hợp của phân tích nhân tố. Trị số này cho biết độ lớn của mẫu được sử dụng có thể cho kết quả phân tích nhân tố đáng tin cậy hay không. KMO có giá trị nằm trong khoảng từ 0 đến 1, với các giá trị cao hơn cho thấy khả năng lấy mẫu đầy đủ tốt hơn. Giá trị KMO từ 0,5 trở xuống cho biết tập dữ liệu không phù hợp để phân tích nhân tố, trong khi giá trị KMO từ 0,6 đến 0,7 cho biết tập dữ liệu là cận biên (marginal) và giá trị KMO từ 0,8 trở lên cho biết tập dữ liệu phù hợp cho phân tích nhân tố.

Công thức của *Trị số KMO*

$$MO_j = \frac{\sum_{i \neq j} r_{ij}^2}{\sum_{i \neq j} r_{ij}^2 + \sum_{i \neq j} u_i}$$

Trong đó:

- $\sum_{i \neq j} r_{ij}^2$ : Tổng các giá trị hiệp phương sai giữa các biến
- $\sum_{i \neq j} u_i$ : Tổng các giá trị phương sai không chéo

Tóm lại, kiểm định KMO là một thước đo thống kê đánh giá mức độ phù hợp của cỡ mẫu để tiến hành phân tích khám phá nhân tố. Giá trị KMO cao cho biết tập dữ liệu phù hợp để phân tích nhân tố, trong khi giá trị KMO thấp cho biết tập dữ liệu không phù hợp để phân tích nhân tố.

- *Kiểm định Bartlett*, là một phương pháp kiểm tra giả thuyết về sự đồng nhất của các phương sai trong một mẫu dữ liệu. Phương pháp này được sử dụng để kiểm tra xem liệu có sự khác biệt đáng kể giữa các phương sai của các biến trong một mẫu dữ liệu hay không. Nếu kiểm định này có ý nghĩa trong thống kê ( $\text{Sig} \leq 0.05$  or  $0.01$ ) thì các biến quan sát có tương quan với nhau trong tổng thể. Ta có thể bác bỏ giả thuyết về tính đồng nhất của các phương sai. ( $H_0$  : Các phương sai của các biến trong mẫu dữ liệu là bằng nhau).
- Trong Phân tích nhân tố khám phá (EFA), Eigenvalue (Trị riêng) được sử dụng để xác định số lượng nhân tố giải thích tốt nhất phương sai trong một tập hợp các biến quan sát. Chỉ những nhân tố có Eigenvalue lớn hơn 1 thì mới được giữ lại trong mô hình. Những nhân tố có Eigenvalue nhỏ hơn 1 sẽ không có tác dụng tóm tắt thông tin tốt hơn một biến gốc.
- *Factor loading (FL)*: Trong Phân tích nhân tố khám phá (EFA), *Factor loading* (trọng số nhân tố) đề cập đến mối tương quan giữa biến quan sát và nhân tố tiềm ẩn. Nó đo lường mức độ mà biến quan sát có liên quan đến

nhân tố tiềm ẩn và nó là một đầu ra quan trọng của phân tích. Là chỉ tiêu để đảm bảo mức ý nghĩa thiết thực của EFA, việc lựa chọn giá trị của EFA phụ thuộc vào cỡ mẫu quan sát và mục đích của nghiên cứu. Factor loading (FL) thường được biểu thị bằng một số từ -1 đến 1, với các giá trị càng gần 1 cho biết mối quan hệ chặt chẽ hơn giữa biến và nhân tố. Factor loading (FL) dương chỉ ra rằng biến có mối liên hệ tích cực với nhân tố, trong khi factor loading (FL) cho thấy mối liên hệ tiêu cực. Factor loading (FL) gần bằng 0 cho thấy rằng biến không liên quan nhiều đến nhân tố. Nếu  $FL \geq 0.3$  là đạt được mức tối thiểu với cỡ mẫu khoảng 300,  $FL \geq 0.4$  được xem là quan trọng và  $FL \geq 0.5$  được xem như là có ý nghĩa thực tiễn. Khi cỡ mẫu khoảng 100 thì nên chọn  $FL \geq 0.55$ , còn nếu cỡ mẫu 50 thì nên chọn  $FL \geq 0.75$ .

- *Component matrix (Rotated component matrix)*: Một phần quan trọng trong bảng kết quả phân tích nhân tố là ma trận nhân tố (component matrix) hay ma trận nhân tố khi các nhân tố được xoay (rotated component matrix). Ma trận nhân tố chứa các hệ số biểu diễn các biến chuẩn hóa bằng các nhân tố (mỗi biến là một đa thức của các nhân tố). Những hệ số tải nhân tố (factor loading) biểu diễn tương quan giữa các biến và các nhân tố. Hệ số này cho biết nhân tố và biến có liên quan chặt chẽ với nhau. Nghiên cứu sử dụng phương pháp trích nhân tố principal components nên các hệ số tải nhân tố phải có trọng số lớn hơn 0.5 thì mới đạt yêu cầu.

Kết thúc bước EFA, chúng ta sẽ tiến hành **tạo nhân tố đại diện (biến đại diện)** phục vụ cho bước phân tích tương quan quan, hồi quy, phương sai. chúng ta thực hiện kiểm định độ tin cậy Cronbach's Alpha và phân tích nhân tố EFA để chọn lọc giữ lại các biến quan sát tốt và loại bỏ các biến quan sát kém chất lượng. Kết thúc bước chọn lọc, chúng ta có được các nhân tố phù hợp nhất với các biến quan sát tốt nhất. Lúc này, chúng ta sẽ cần chuyển hướng việc đo lường biến quan sát về

đo lường nhân tố để kết luận được các giả thuyết đặt ra, bước chuyển đổi này được gọi là tạo nhân tố đại diện.

- Tạo nhân tố đại diện bằng trung bình cộng (mean)
- Tạo nhân tố đại diện bằng tổng giá trị (sum)
- Tạo nhân tố đại diện bằng điểm nhân tố.

## **1.2.6 Phân tích hồi quy đa biến**

### **1.2.6.1 Phân tích tương quan**

Phân tích tương quan là quá trình đo lường mức độ liên quan giữa hai hoặc nhiều biến trong một bộ dữ liệu. Nó cho phép nhà nghiên cứu đánh giá mối quan hệ giữa các biến, từ đó giúp hiểu rõ hơn về mối liên hệ giữa các yếu tố trong nghiên cứu.

Phương pháp phân tích này sử dụng hệ số Pearson để lượng hóa mức độ chặt chẽ về mối quan hệ giữa hai biến định lượng. Giá trị tuyệt đối của hệ số Pearson càng gần đến 1 thì chứng tỏ hai biến có mối liên hệ tương quan tuyến tính càng chặt chẽ.

$r < \pm 0.20$ : Không có mối quan hệ tương quan

$0.21 \leq r \leq 0.40$ : Tương quan yếu

$0.41 \leq r \leq 0.60$ : Tương quan trung bình

$0.61 \leq r \leq 0.80$ : Tương quan mạnh

$0.81 \leq r \leq 1.00$ : Tương quan rất mạnh

### **1.2.6.2 Phân tích hồi quy đa biến**

Phân tích hồi quy đa biến là phương pháp phân tích đa biến trong thống kê, sử dụng để xác định mối quan hệ giữa một biến phụ thuộc và nhiều biến độc lập. Nó cho phép nhà nghiên cứu xác định ảnh hưởng của nhiều yếu tố độc lập đến một biến phụ thuộc và đưa ra dự đoán cho giá trị của biến phụ thuộc dựa trên các giá trị của các biến độc lập.

Phân tích hồi quy đa biến có thể được thực hiện bằng cách sử dụng các phương pháp như phân tích hồi quy tuyến tính đa biến và phân tích hồi quy logistic đa biến.

Trong phân tích hồi quy tuyến tính đa biến, ta sử dụng một mô hình tuyến tính để mô tả mối quan hệ giữa biến phụ thuộc và các biến độc lập, trong khi trong phân tích hồi quy logistic đa biến, ta sử dụng một mô hình logistic để mô tả mối quan hệ giữa biến phụ thuộc và các biến độc lập. Quá trình phân tích và đi đến kết luận của phân tích mô hình đa biến chúng ta cần lưu ý những điểm sau:

***a. Kiểm tra hệ số phóng đại phương sai***

Hệ số phóng đại phương sai (Variance Inflation Factor - VIF) được sử dụng để kiểm tra độ tương quan giữa các biến độc lập trong một mô hình hồi quy tuyến tính. VIF đo lường mức độ mà phương sai của một biến độc lập được giải thích bởi các biến độc lập khác trong mô hình.

Một VIF cao cho thấy rằng biến đó có mức độ tương quan cao với các biến độc lập khác và có thể gây ra vấn đề về đa cộng tuyến trong mô hình. Một giá trị VIF dưới 5 thường được coi là chấp nhận được trong phân tích hồi quy tuyến tính.

Nếu giá trị này lớn hơn 5 thì có sự xuất hiện đa cộng tuyến (Multicollinearity) trong mô hình. Đa cộng tuyến là hiện tượng tạo nên từ mối quan hệ tương quan mạnh giữa các biến độc lập với nhau trong mô hình hồi quy tuyến tính. Hiện tượng này được thể hiện dưới dạng hàm số sau khi vi phạm giả thuyết của mô hình hồi quy tuyến tính cổ điển. (Giả thuyết vi phạm: Các biến độc lập không có quan hệ tuyến tính với nhau)

***b. Hệ số  $R^2$  đã được điều chỉnh (adjusted  $R$  square)***

Hệ số  $R^2$  được sử dụng để đo lường mức độ giải thích của một mô hình hồi quy tuyến tính đối với dữ liệu. Nó là tỷ lệ phần trăm của sự biến thiên của biến phụ thuộc được giải thích bởi mô hình hồi quy tuyến tính. Hệ số  $R^2$  đã được điều chỉnh (adjusted  $R^2$ ) là một biến thể của hệ số  $R^2$ , nó được sử dụng để cân nhắc đến số lượng biến độc lập được sử dụng trong mô hình hồi quy tuyến tính. Nó được tính bằng cách

điều chỉnh hệ số  $R^2$  bằng cách thêm vào một điều chỉnh cho số lượng biến độc lập trong mô hình và số lượng mẫu trong tập dữ liệu.

$$R\text{-squared} = R^2 = \text{SSE}/\text{SST} = 1 - \text{SSR}/\text{SST}$$

Total sum of squares (SST) = explained sum of squares (SSE) + residual sum of squares (SSR)

R-squared đo tỷ lệ của tổng biến thể được giải thích bằng hồi quy. R-squared là một thước đo phù hợp tốt (good fit).

Công thức tính hệ số  $R^2$  đã được điều chỉnh như sau:

$$\text{Adjusted } R^2 = 1 - [(1 - R^2) * (n - 1) / (n - k - 1)]$$

Trong đó:

- $R^2$  là hệ số  $R^2$
- $n$  là số lượng mẫu trong tập dữ liệu
- $k$  là số lượng biến độc lập trong mô hình.

Hệ số  $R^2$  đã được điều chỉnh thường có giá trị thấp hơn so với hệ số  $R^2$  vì nó tính toán mức độ giải thích của mô hình dựa trên số lượng biến độc lập trong mô hình và số lượng mẫu trong tập dữ liệu.

***Ví dụ minh họa:*** Mô hình kinh tế lượng mô tả các yếu tố quyết định tiền lương của người lao động.

$$\widehat{\text{Wage}} = \hat{\beta}_0 + \hat{\beta}_1 \text{education} + \hat{\beta}_2 \text{experience} + \hat{\beta}_3 \text{tenure}$$

Kết quả phân tích hồi quy tuyến tính bội in R

| Source   | SS         | df  | MS         | Number of obs | = | 526    |
|----------|------------|-----|------------|---------------|---|--------|
| Model    | 2194.11162 | 3   | 731.370541 | F(3, 522)     | = | 76.87  |
| Residual | 4966.30269 | 522 | 9.51398982 | Prob > F      | = | 0.0000 |
|          |            |     |            | R-squared     | = | 0.3064 |
|          |            |     |            | Adj R-squared | = | 0.3024 |
| Total    | 7160.41431 | 525 | 13.6388844 | Root MSE      | = | 3.0845 |

| wage   | Coef.     | Std. Err. | t     | P> t  | [95% Conf. Interval] |           |
|--------|-----------|-----------|-------|-------|----------------------|-----------|
| educ   | .5989651  | .0512835  | 11.68 | 0.000 | .4982176             | .6997126  |
| exper  | .0223395  | .0120568  | 1.85  | 0.064 | -.0013464            | .0460254  |
| tenure | .1692687  | .0216446  | 7.82  | 0.000 | .1267474             | .2117899  |
| _cons  | -2.872735 | .7289643  | -3.94 | 0.000 | -4.304799            | -1.440671 |

R-squared = SS Model /SS Total = 2194 / 7160 = 1 – SSR/SST = 1 – 4966/7160 = 0.306  
 Adj R-squared = 1 – [SSR/(n-k-1)] / [SST/(n-1)] = 1 – (4966/522)/ (7160/525) = 0.302  
 30% of the variation in wage is explained by the regression and the rest is due to error.

Như vậy, hệ số ***adjusted R<sup>2</sup>*** bằng 0.3 nghĩa là mô hình hồi qui tuyến tính bội đã xây dựng phù hợp với tập dữ liệu là 30%. Nói cách khác, khoảng 30 % sự thay đổi tiền lương quan sát có thể được giải thích bởi sự khác biệt của 03 thành phần: giáo dục, kinh nghiệm và thâm niên.

**c. Kiểm định giả thuyết:** Kiểm định giả thuyết về độ phù hợp của mô hình. Kiểm định giả thuyết về ý nghĩa của hệ số hồi quy  $\beta$ . Xác định mức độ ảnh hưởng của các nhân tố đến biến phụ thuộc. Kết quả phân tích dữ liệu sẽ cho ta biết liệu giả thuyết có được chấp nhận hay không. Nếu giả thuyết không được chấp nhận, ta sẽ phải đưa ra giải thích cho kết quả đó và cân nhắc đến các giả thuyết thay thế.

Tuy nhiên, cần lưu ý rằng kiểm định giả thuyết chỉ có thể đưa ra kết luận tạm thời về tính đúng đắn của giả thuyết. Việc kiểm định giả thuyết cũng không đảm bảo rằng giả thuyết là chính xác 100%.

**d. Hiện tượng tự tương quan:** Tự tương quan được hiểu như là sự tương quan giữa các thành phần của dãy số thời gian hoặc không gian. Nguyên nhân của hiện tượng này có thể là khách quan (do quán tính, do độ trễ của dữ liệu) hoặc có thể là chủ quan (do chọn mô hình sai, do quá trình xử lý số liệu). Hậu quả của hiện

tượng này như sau: các ước lượng bình phương bé nhất vẫn là ước lượng tuyến tính, không chệch nhưng không phải là ước lượng hiệu quả hoặc là phương sai ước lượng được của các ước lượng bình phương bé nhất thường bị chệch, các kiểm định T và F không đáng tin cậy.

Để giải quyết hiện tượng này chúng ta sử dụng phương pháp thống kê Durbin – Watson. Phương pháp thống kê Durbin-Watson là một phương pháp thống kê được sử dụng để kiểm tra một số giả định cơ bản trong phân tích hồi quy tuyến tính, bao gồm sự độc lập của các sai số và sự tương quan không đồng nhất của các sai số.

Thống kê kiểm định Durbin-Watson, thường ký hiệu là  $d$ , được tính như sau:

$$d = \frac{\sum_{t=2}^T (e_t - e_{t-1})^2}{\sum_{t=1}^T e_t^2}$$

Trong đó:

- $T$ : The total number of observations
- $e_t$ : The  $t^{\text{th}}$  residual from the regression model

Hệ số này có giá trị từ 0 đến 4 và được tính dựa trên sự khác biệt giữa các sai số liên tiếp. Nếu giá trị của hệ số Durbin-Watson gần bằng 2, tức là không có tương quan tức thời giữa các sai số liên tiếp, thì giả định về sự độc lập của các sai số được chấp nhận. Tuy nhiên, nếu giá trị của hệ số Durbin-Watson thấp hơn 2, thì có sự tương quan giữa các sai số liên tiếp và giả định về sự độc lập của các sai số không được chấp nhận. Ngược lại, nếu giá trị của hệ số Durbin-Watson cao hơn 2, thì có sự tương quan âm giữa các sai số liên tiếp, tức là sai số của mô hình có

xu hướng thay đổi giữa các quan sát. Trong trường hợp này, các phương pháp khác như phân tích chuỗi thời gian có thể được sử dụng để xử lý sự tương quan này.

Trong bảng giá trị tới hạn Durbin-Watson, vùng tới hạn nằm trong khoảng từ 0,97 (dL) đến 1,33 (dU) đối với  $N = 12$  với ý nghĩa 5%. Vì thống kê thử nghiệm Durbin-Watson ( $DW = 2,58$ ) cao hơn 1,33 ( $DW > dU$ ), chúng ta không bác bỏ giả thuyết null ( $p > 0,05$ ) rằng không có tự tương quan. Vì giá trị p thu được từ thử nghiệm Durbin-Watson không đáng kể ( $d = 2,584$ ,  $p = 0,470$ ), chúng ta không bác bỏ giả thuyết null. Do đó, chúng tôi kết luận rằng phần còn lại không tự tương quan.

### 1.2.6.3 Kiểm định tham số trung bình tổng thể và phân tích ANOVA

#### a. Kiểm định tham số trung bình tổng thể

Phép kiểm định này được sử dụng trong trường hợp chúng ta cần so sánh trị trung bình về một tiêu chí nghiên cứu nào đó giữa hai đối tượng mà chúng ta quan tâm. Trước khi thực hiện kiểm định trung bình, ta cần thực hiện kiểm định sự bằng nhau của hai phương sai tổng thể. *Kiểm định này có tên là Levene (Levene's test)*, với giả thiết  $H_0$  rằng phương sai của hai tổng thể bằng nhau.

Với mức ý nghĩa 0,05, chúng ta có:

- Nếu  $Sig < 0.05$ : phương sai giữa các nhóm không có sự khác biệt.
- Nếu  $Sig \geq 0.05$ : phương sai giữa các nhóm có sự khác biệt.

Kết quả của việc chấp nhận hay bác bỏ  $H_0$  ảnh hưởng quan trọng đến việc chúng ta sẽ lựa chọn loại kiểm định giả thiết về sự bằng nhau của hai trung bình tổng thể nào: kiểm định trung bình với phương sai bằng nhau hay kiểm định trung bình với phương sai khác nhau.

Trong trường hợp  $p\text{-value} < 0.05$  thì giả định  $H_0$  được chấp nhận điều đó có nghĩa là phương sai giữa các nhóm có sự khác biệt. Do đó kết quả ở bảng ANOVA không sử dụng được vì thế phải sử dụng kiểm định Kruskal-Wallis. Còn  $p\text{-value} > 0.05$  thì phương sai giữa các nhóm không sự khác biệt và kết quả ở bảng ANOVA sẽ được sử dụng cho các giá trị này.

Kiểm định Kruskal-Wallis là một kiểm định phi tham số được sử dụng để kiểm tra sự khác biệt về giá trị trung vị giữa các nhóm dữ liệu độc lập. Nếu  $p\text{-value}$  được tính toán bởi kiểm định Kruskal-Wallis là nhỏ hơn một ngưỡng xác định trước (thường là 0,05), thì chúng ta có thể kết luận rằng có sự khác biệt đáng kể giữa các nhóm dữ liệu.

## **b. Phân tích ANOVA**

ANOVA (Analysis of Variance) tests là các phương pháp thống kê được sử dụng để so sánh giá trị trung bình của hai hoặc nhiều nhóm dữ liệu độc lập. Phân tích ANOVA phân tích sự phân tán giữa các nhóm dữ liệu và sử dụng phương sai để đánh giá sự khác biệt giữa các giá trị trung bình.

ANOVA kiểm tra xem có sự khác biệt về giá trị trung bình của các nhóm ở mỗi cấp độ của biến độc lập hay không.

- Giả thuyết không ( $H_0$ ) của ANOVA không có sự khác biệt của các giá trị trung bình.
- Giả thuyết thay thế ( $H_a$ ) là các giá trị trung bình khác nhau.

nếu  $p\text{-value} \leq 0.05$ : bác bỏ  $H_0 \rightarrow$  đủ điều kiện để khẳng định có sự khác biệt giữa các nhóm đối với biến phụ thuộc. Ngược lại,  $p\text{-value} > 0.05$ : chấp nhận  $H_0 \rightarrow$  chưa đủ điều kiện để khẳng định có sự khác biệt giữa các nhóm đối với biến phụ thuộc

Khi có sự khác biệt thì có thể phân tích sâu hơn để tìm ra sự khác biệt như thế nào giữa các nhóm quan sát bằng các kiểm định Tukey, LSD, Bonferroni, Duncan như hình dưới. Kiểm định sâu anova gọi là kiểm định Post-Hoc.

Phân tích ANOVA có ba loại chính: ANOVA một chiều (one-way ANOVA) dùng để so sánh giá trị trung bình của một biến phụ thuộc (outcome variable) trong các nhóm dữ liệu độc lập; ANOVA hai chiều (two-way ANOVA) dùng để xác định sự tương tác giữa hai biến độc lập đối với một biến phụ thuộc; và ANOVA nhiều chiều (n-way ANOVA) được sử dụng để so sánh giá trị trung bình của một biến phụ thuộc trong nhiều nhóm dữ liệu độc lập.

***Ví dụ ANOVA một chiều: Chúng tôi kiểm định ảnh hưởng của 3 loại phân bón đến năng suất cây trồng.***

Kết quả: Phân tích Anova in R

```
one.way <- aov(yield ~ fertilizer, data = crop.data)

summary(one.way)
```

|            | Df | Sum Sq | Mean Sq | F value | Pr(>F)    |
|------------|----|--------|---------|---------|-----------|
| fertilizer | 2  | 6.07   | 3.0340  | 7.863   | 7e-04 *** |
| Residuals  | 93 | 35.89  | 0.3859  |         |           |

---  
Signif. codes: 0 '\*\*\*' 0.001 '\*\*' 0.01 '\*' 0.05 '.' 0.1 ' ' 1

Chú ý các giá trị

- Cột F value là thống kê thử nghiệm từ F-test. Đây là bình phương trung bình của mỗi biến độc lập chia cho bình phương trung bình của phần dư. Giá trị F càng lớn thì càng có nhiều khả năng biến động do biến độc lập gây ra là có thật chứ không phải do ngẫu nhiên.
- Cột Pr (>F) là giá trị p của thống kê F. Điều này cho thấy khả năng giá trị F được tính toán từ thử nghiệm sẽ xảy ra như thế nào nếu giả thuyết không về không có sự khác biệt giữa các phương tiện nhóm là đúng.

Kết quả cho thấy rằng giá trị  $p$  của biến phân bón thấp  $p.value = 7e - 4$  ( $p < 0,001$ ), vì vậy có vẻ như loại phân bón được sử dụng có tác động thực sự đến năng suất cây trồng cuối cùng. Có nghĩa là có sự khác nhau giữa các loại phân bón tác động đến năng suất

## **PHẦN II. DỰ ÁN NGHIÊN CỨU : CẤU TRÚC ĐỀ CƯƠNG CHUNG CỦA MỘT DỰ ÁN NGHIÊN CỨU**

**LỜI CAM DOAN**

**LỜI CẢM ƠN**

**TÓM TẮT LUẬN VĂN**

**MỤC LỤC**

**DANH MỤC HÌNH**

**DANH MỤC BẢNG BIỂU**

**DANH MỤC CÁC TỪ VIẾT TẮT**

### **CHƯƠNG I: TỔNG QUAN VỀ NGHIÊN CỨU (10 GIỜ LÀM VIỆC)**

**1.1 XÁC ĐỊNH ĐỀ TÀI NGHIÊN CỨU**

**1.2 MỤC TIÊU NGHIÊN CỨU**

**1.3 CÂU HỎI NGHIÊN CỨU**

**1.4 ĐỐI TƯỢNG VÀ PHẠM VI NGHIÊN CỨU**

**1.6 Ý NGHĨA CỦA NGHIÊN CỨU**

**1.7 KẾT CẤU BÁO CÁO NGHIÊN CỨU**

### **CHƯƠNG 2: CƠ SỞ LÝ THUYẾT VÀ MÔ HÌNH NGHIÊN CỨU (10 GIỜ LÀM VIỆC)**

**2.1. CƠ SỞ LÝ THUYẾT**

2.1.1.

2.1.2.

*+ Sinh viên cần phải làm rõ các thuật ngữ liên quan đến chủ đề mà bạn chọn, các khái niệm và các nhận định phải có cơ sở về mặt khoa học.*

## **2.2. MÔ HÌNH NGHIÊN CỨU LÝ THUYẾT & THỰC TIỄN**

### **2.2.1. MÔ HÌNH NGHIÊN CỨU LÝ THUYẾT**

2.2.1.1. (Công trình Lý thuyết khoa học 1)

2.2.1.2. (Công trình Lý thuyết khoa học 2)

2.2.2.3. (Công trình Lý thuyết khoa học 3)

...vv.

*+ Các lý thuyết phải có sự liên quan và làm nền tảng cho chủ đề nghiên cứu.*

### **2.2.2. MÔ HÌNH NGHIÊN CỨU THỰC TIỄN**

2.2.2.1. (Công trình bài báo khoa học thực tiễn 1)

2.2.2.2. (Công trình bài báo khoa học thực tiễn 2)

2.2.2.3. (Công trình bài báo khoa học thực tiễn 3)

..vv

*+ Cần kết hợp bài báo thực tiễn trong nước và trong nước. Các mô hình nghiên cứu thực tiễn có liên quan đến tên đề tài mà các bạn chọn .*

## **2.3. MÔ HÌNH NGHIÊN CỨU ĐỀ XUẤT**

### **2.3.1. XÂY DỰNG CÁC GIẢ THUYẾT NGHIÊN CỨU**

*+ Sinh viên xây dựng giả thuyết nghiên cứu dựa những lập luận khoa học trước đây để xây dựng giả thuyết nghiên cứu của mình cho các biến độc lập và biến phụ thuộc trong mô hình nghiên cứu đề xuất của mình.*

### **2.3.2. MÔ HÌNH NGHIÊN CỨU ĐỀ XUẤT**

*+ Mô hình đề xuất xây dựng cần chú ý dựa trên những cơ sở lý thuyết ở mục 2.2.1 và 2.2.2 (các biến).*

## **CHƯƠNG 3. PHƯƠNG PHÁP NGHIÊN CỨU (10 GIỜ LÀM VIỆC)**

### **3.1. THIẾT KẾ NGHIÊN CỨU**

#### 3.1.1. Phương pháp nghiên cứu

#### 3.1.2. Quy trình nghiên cứu

#### 3.1.3. Nghiên cứu định tính

##### 3.1.3.1. Thiết kế nghiên cứu định tính

##### 3.1.3.2. Kết quả nghiên cứu định tính

#### 3.1.4. Nghiên cứu định lượng

##### 3.1.4.1. Phương pháp chọn mẫu và cỡ mẫu

##### 3.1.4.2. Cách thức thu thập và xử lý dữ liệu

### **3.2. THIẾT KẾ THANG ĐO**

*Cách thiết kế thang đo chú ý hai điểm sau:*

- *Xác định đó là thang đo nguyên nhân hay kết quả*

- *Các thang đo cần có sự thừa kế trên cơ sở các thang đo của các lý thuyết khoa học hay các tác giả đã được công nhận*
- *Các thang đo tránh trùng lặp ý nghĩa nhau*

## **CHƯƠNG 4 KẾT QUẢ NGHIÊN CỨU (10 GIỜ LÀM VIỆC)**

### **4.1. PHÂN TÍCH MÔ TẢ THỐNG KÊ**

#### 4.1.1. Mô tả mẫu

#### 4.1.2. Thống kê mô tả các biến nghiên cứu

### **4.2. KIỂM ĐỊNH ĐỘ TIN CẬY CRONBACH'S ALPHA CỦA THANG ĐO**

#### 4.2.1. Thang đo biến độc lập

#### 4.2.2. Thang đo biến phụ thuộc

### **4.3. PHÂN TÍCH NHÂN TỐ KHÁM PHÁ (EFA)**

### **4.4. PHÂN TÍCH HỒI QUY**

#### 4.4.1. Phân tích ma trận tương quan

#### 4.4.2. Phân tích hồi quy

##### 4.4.2.1 Đánh giá độ phù hợp của mô hình

*Chú ý: các bạn đọc chi tiết phần I và phần III của tài liệu này.*

- ***Nhận biết phân phối chuẩn (Normal distribution)***
  - + *Histograms with normal curve (or)*

- + *Normal Q- Q Plot*
- + *OR Kiểm định Kolmogorov – smirnov khi cỡ mẫu lớn hơn 50 hoặc phép kiểm Shapiro – wilk*
- **Đa cộng tuyến (Multicollinearity)**
  - + *Variance Inflating factor (VIF)*
  - + *Giá trị dung sai tolerance*
- **Phương sai thay đổi (heteroscedasticity)**
  - + *White, the Breusch-Pagan test, Glesjer, tương quan hạng Spearman (or)*
  - + *Biểu đồ Scatter Plot*
- **Kết quả phân tích hồi quy đa tuyến**
  - + *F-Test*
  - + *Regression Coefficient*
  - + *Coefficient of Determination ( $R^2$ , adjust  $R^2$ )*
  - + *T-Test*

4.4.2.2. Kết quả kiểm định các giả thuyết nghiên cứu trong mô

## 4.5 KIỂM ĐỊNH SỰ KHÁC BIỆT CỦA CÁC BIẾT KIỂM SOÁT

*Giới tính, tình trạng hôn nhân, số người nuôi dưỡng, tuổi, nghề nghiệp, trình độ giáo dục, ...vv.*

## CHƯƠNG 5: KẾT LUẬN VÀ KIẾN NGHỊ (5 GIỜ LÀM VIỆC)

### 5.1. KẾT QUẢ NGHIÊN CỨU VÀ NHỮNG ĐÓNG GÓP CỦA NGHIÊN CỨU

### 5.2. HÀM Ý CỦA KẾT QUẢ NGHIÊN CỨU

### **5.3. HẠN CHẾ VÀ HƯỚNG NGHIÊN CỨU TIẾP THEO**

#### **TÀI LIỆU THAM KHẢO**

#### **PHỤ LỤC**

# PHẦN III: HƯỚNG DẪN ĐỌC Ý NGHĨA VÀ CÁC BƯỚC THỰC HIỆN NCKH

## 1. PHƯƠNG PHÁP HỒI QUY ĐA BIẾN SỬ DỤNG TRÊN SPSS

### Descriptive Statistics

- Mean, min, max, Frequency, Percent

### Reliability analysis

- Cronbach's Alpha

### Exploratory Factor

- KMO, Bartlett test
- Trị số Eigenvalue (SS loadings (Eigenvalue))
- Tổng phương sai trích (Total Variance Explained)
- Hệ số tải nhân tố (Factor Loading)

### Pearson correlation

- Kỳ vọng: sig. tương quan giữa độc lập với phụ thuộc nhỏ hơn 0.05 và hệ số tương quan càng cao càng tốt.
- Kỳ vọng: (1) sig. tương quan giữa các biến độc lập lớn hơn 0.05 hoặc (2) sig nhỏ hơn 0.05 và hệ số tương quan sẽ càng thấp càng tốt (nên dưới 0.5).

### Regression analysis

- **Nhận biết phân phối chuẩn (Normal distribution)**
  - + Histograms with normal curve (or)
  - + Normal Q- Q Plot
  - + Kiểm định Kolmogorov – smirnov khi cỡ mẫu lớn hơn 50 hoặc phép kiểm Shapiro – wilk
- **Đa cộng tuyến (Multicollinearity)**
  - + Variance Inflating factor (VIF)
  - + Giá trị dung sai tolerance
- **Phương sai thay đổi (heteroscedasticity)**
  - + White, the Breusch-Pagan test, Glesjer, tương quan hạng Spearman (or)
  - + Biểu đồ Scatter Plot
- **Kết quả phân tích hồi quy đa tuyến**
  - + F-Test
  - + Regression Coefficient
  - + Coefficient of Determination ( $R^2$ , adjust  $R^2$ )
  - + T-Test
- **Kiểm định sự khác biệt về tổng thể**
  - + Levene (Levene's test)
  - + ANOVA

## 1.1: Thống kê mô tả (Descriptive Statistics)

Thống kê mô tả là các hệ số thông tin ngắn gọn tóm tắt một tập dữ liệu nhất định, có thể là đại diện cho toàn bộ tổng thể hoặc một mẫu của tổng thể.

| Lí thuyết                                                        | Hàm R     |
|------------------------------------------------------------------|-----------|
| Số trung bình: $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$          | mean (x)  |
| Phương sai: $s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$ | var (x)   |
| Độ lệch chuẩn: $s = \sqrt{s^2}$                                  | sd (x)    |
| Sai số chuẩn (standard error): $SE = \frac{s}{\sqrt{n}}$         | Không có  |
| Trị số thấp nhất                                                 | min (x)   |
| Trị số cao nhất                                                  | max (x)   |
| Toàn cự (range)                                                  | range (x) |

### Chú ý :

Độ lệch chuẩn càng cao thì hiệu thị sự chênh lệch càng nhiều trong đáp án của người khảo sát.

### Interquartile Range

- The upper fence =  $Q3 + IQR * 1.5$
- The lower fence =  $Q1 - IQR * 1.5$
- $IQR = Q3 - Q1$

Thống kê các biến kiểm soát trong nghiên cứu với Frequency, Percent

## 1.2. Phân tích độ tin cậy (Reliability analysis)

### 1.2.1. Cronbach's Alpha

Cronbach's Alpha là một cách để đo lường tính nhất quán bên trong của một bảng câu hỏi hoặc khảo sát.

### Các hệ số

- Cronbach's Alpha: Hệ số Cronbach's Alpha

- N of Items: Số lượng biến quan sát
- Scale Mean if Item Deleted: Trung bình thang đo nếu loại biến
- Scale Variance if Item Deleted: Phương sai thang đo nếu loại biến

biến

- Corrected Item-Total Correlation: Tương quan biến tổng
- Cronbach's Alpha if Item Deleted: Hệ số Cronbach's Alpha nếu loại biến

loại biến

**Cronbach's Alpha nằm trong khoảng từ 0 đến 1, với các giá trị càng cao cho thấy khảo sát hoặc bảng câu hỏi đáng tin cậy hơn.**

Bảng sau đây mô tả các giá trị khác nhau của Cronbach's Alpha thường được diễn giải như sau:

| Cronbach's Alpha        | Internal consistency |
|-------------------------|----------------------|
| $0.9 \leq \alpha$       | Excellent            |
| $0.8 \leq \alpha < 0.9$ | Good                 |
| $0.7 \leq \alpha < 0.8$ | Acceptable           |
| $0.6 \leq \alpha < 0.7$ | Questionable         |
| $0.5 \leq \alpha < 0.6$ | Poor                 |
| $\alpha < 0.5$          | Unacceptable         |

Theo Nunnally (1978), một thang đo tốt nên có độ tin cậy Cronbach's Alpha từ 0.7 trở lên. Hair và cộng sự (2009) cũng cho rằng, một thang đo đảm bảo tính đơn hướng và đạt độ tin cậy nên đạt ngưỡng Cronbach's Alpha từ 0.7 trở lên, tuy nhiên, với tính chất là một nghiên cứu khám phá sơ bộ, ngưỡng Cronbach's Alpha là 0.6 có thể chấp nhận được. Hệ số Cronbach's Alpha càng cao thể hiện độ tin cậy của thang đo càng cao.

**Corrected Item - Total Correlation** thể hiện mối quan hệ giữa biến quan sát đó với tất cả các biến quan sát còn lại trong cùng 1 thang đo (cùng 1 nhóm).

**Các ngưỡng cần chú ý:**

- 0.30 by Nunnally, J. C., & Bernstein, I. H. (1994)
- 0.30 by Cristobal et al. (2007)
- 0.40 by Loiacono et al. (2002)
- 0.50 by Francis and White (2002) and Kim and Stoel (2004)

**Chú ý\***

- Nếu thang đo có hệ số Cronbach's Alpha dưới 0.6, chúng ta **chưa vội** kết luận thang đo không có độ tin cậy mà cần kiểm tra và loại hết biến quan sát có giá trị Cronbach's Alpha if Item Deleted cao hơn mức 0.6. Đến khi loại hết rồi mà hệ số Cronbach's Alpha vẫn dưới 0.6 thì mới kết luận.
- Nếu hệ số Cronbach's Alpha của nhóm đã **đủ tiêu** chuẩn thì việc xuất hiện biến quan sát có Cronbach's Alpha if Item Deleted lớn hơn Cronbach's Alpha của nhóm nhưng tương quan biến tổng lớn hơn 0.3 thì chúng ta không cần loại biến quan sát đó đi.
- Nếu hệ số Cronbach's Alpha của nhóm **chưa đủ tiêu** chuẩn thì việc xuất hiện biến quan sát có Cronbach's Alpha if Item Deleted lớn hơn Cronbach's Alpha của nhóm nhưng tương quan biến tổng lớn hơn 0.3 thì chúng ta nên loại biến quan sát đó đi để cải thiện độ tin cậy thang đo cho tới khi hệ số Cronbach's Alpha của nhóm đạt tiêu chuẩn.
- Nếu hệ số Cronbach's Alpha của nhóm **chưa đủ tiêu chuẩn**, chúng ta đã loại các biến quan sát có Cronbach's Alpha if Item Deleted lớn hơn Cronbach's Alpha của nhóm nhưng thang đo vẫn không đủ tiêu chuẩn. Khi đó, thang đo không đảm bảo độ tin cậy cho nghiên cứu, cần loại bỏ cả thang đo này.

- Nếu sự chênh lệch giữa Cronbach's Alpha của nhóm với Cronbach's Alpha if Item Deleted của biến quan sát là đáng kể từ **0.3 trở lên**. Chúng ta sẽ loại biến quan sát đó để tăng thêm độ tin cậy của thang đo.

### 1.3. Phân tích nhân tố khám phá

**Phân tích nhân tố khám phá**, gọi tắt là EFA, dùng để rút gọn một tập hợp k biến quan sát thành một tập F (với  $F < k$ ) các nhân tố có ý nghĩa hơn. Với kiểm định độ tin cậy thang đo Cronbach Alpha, chúng ta đang đánh giá mối quan hệ giữa các biến trong cùng một nhóm, cùng một nhân tố, chứ không xem xét mối quan hệ giữa tất cả các biến quan sát ở các nhân tố khác. Trong khi đó, EFA xem xét mối quan hệ giữa các biến ở tất cả các nhóm (các nhân tố) khác nhau nhằm phát hiện ra những biến quan sát tải lên nhiều nhân tố hoặc các biến quan sát bị phân sai nhân tố từ ban đầu.

#### Các tiêu chí trong phân tích EFA

- **Hệ số KMO (Kaiser-Meyer-Olkin)** là một chỉ số dùng để xem xét sự thích hợp của phân tích nhân tố. Trị số của KMO phải đạt giá trị 0.5 trở lên ( **$0.5 \leq KMO \leq 1$** ) là điều kiện đủ để phân tích nhân tố là phù hợp. Nếu trị số này nhỏ hơn 0.5, thì phân tích nhân tố có khả năng không thích hợp với tập dữ liệu nghiên cứu.
- **Kiểm định Bartlett (Bartlett's test of sphericity)** dùng để xem xét các biến quan sát trong nhân tố có tương quan với nhau hay không. Kiểm định Bartlett có ý nghĩa thống kê (**sig Bartlett's Test < 0.05**), chứng tỏ các biến quan sát có tương quan với nhau trong nhân tố.
- **Trị số Eigenvalue** là một tiêu chí sử dụng phổ biến để xác định số lượng nhân tố trong phân tích EFA. Với tiêu chí này, chỉ có những nhân tố nào có **Eigenvalue > 1** mới được giữ lại trong mô hình phân tích.
- **Tổng phương sai trích (Total Variance Explained)  $\geq 50\%$**  cho thấy mô hình EFA là phù hợp. Coi biến thiên là 100% thì trị số này thể hiện các nhân

tổ được trích cô đọng được bao nhiêu % và bị thất thoát bao nhiêu % của các biến quan sát.

- **Hệ số tải nhân tố (Factor Loading)** hay còn gọi là trọng số nhân tố, giá trị này biểu thị mối quan hệ tương quan giữa biến quan sát với nhân tố.

Theo *Hair và cộng sự (2010), Multivariate Data Analysis* hệ số tải từ 0.5 là biến quan sát đạt chất lượng tốt, tối thiểu nên là 0.3.

Hair và cộng sự cũng cho rằng, giá trị tiêu chuẩn của hệ số tải Factor Loading nên được xem xét cùng kích thước mẫu.

| <b>Giá trị Factor Loading</b> | <b>Kích thước mẫu tối thiểu có ý nghĩa thống kê</b> |
|-------------------------------|-----------------------------------------------------|
| <b>0.3</b>                    | <b>350</b>                                          |
| <b>0.35</b>                   | <b>250</b>                                          |
| <b>0.40</b>                   | <b>200</b>                                          |
| <b>0.45</b>                   | <b>150</b>                                          |
| <b>0.50</b>                   | <b>120</b>                                          |
| <b>0.55</b>                   | <b>100</b>                                          |
| <b>0.60</b>                   | <b>85</b>                                           |
| <b>0.65</b>                   | <b>70</b>                                           |
| <b>0.70</b>                   | <b>60</b>                                           |
| <b>0.75</b>                   | <b>50</b>                                           |

### **Các dạng biến xấu trong EFA**

| Rotated component matrix <sup>@</sup> |      |      |      |      |      |
|---------------------------------------|------|------|------|------|------|
| Component                             |      |      |      |      |      |
|                                       | 1    | 2    | 3    | 4    | 5    |
| LD1                                   | .825 |      |      |      |      |
| LD2                                   | .800 |      |      |      |      |
| LD3                                   | .755 |      |      |      |      |
| DT3                                   | .640 | .523 |      |      |      |
| DT1                                   |      | .822 |      |      |      |
| DT4                                   |      | .796 |      |      |      |
| DT2                                   |      | .748 |      |      |      |
| DN1                                   |      |      | .836 |      |      |
| DN3                                   |      |      | .825 |      |      |
| DN2                                   |      |      | .796 |      |      |
| TL4                                   |      |      |      | .769 |      |
| TL3                                   |      |      |      | .760 |      |
| TL1                                   |      |      |      | .749 |      |
| DN4                                   |      |      |      |      | .899 |
| TL2                                   |      |      |      | .421 |      |

#### **Trường hợp 1: Biến quan sát không đảm bảo hệ số tải tiêu chuẩn**

Tác giả chọn ngưỡng hệ số tải là 0.5 và quy định cho ma trận xoay chỉ hiển thị các biến quan sát có hệ số tải từ 0.4 trở lên (các hệ số tải nhỏ hơn 0.4 sẽ bị ẩn đi). Biến TL2 trong bảng Rotated Component Matrix có hệ số tải cao nhất là  $0.421 < 0.5$  (hệ số tải biến TL2 ở các ô còn lại dưới 0.4 nên đã bị ẩn đi), biến này sẽ được loại bỏ.

#### **Trường hợp 2: Biến quan sát chỉ nằm tách biệt một mình ở một nhân tố**

Trong bảng ma trận xoay ở ví dụ trên, biến DN4 nằm tách biệt một mình ở nhân tố số 5. Chúng ta nên loại bỏ biến quan sát này đi.

### **Trường hợp 3: Tải lên nhiều nhóm nhân tố và chênh lệch hệ số tải nhỏ hơn 0.2**

Matt C. Howard (2015) cho rằng, nếu một biến quan sát tải lên ở hai nhân tố nhưng chênh lệch hệ số tải nhỏ hơn 0.2, biến quan sát nên được xem xét loại bỏ.

Trong bảng ma trận xoay ở ví dụ trên, biến DT3 tải lên hai nhân tố số 1 và số 2 với hệ số tải lần lượt ở hai nhân tố là 0.640 và 0.523. Hiệu số hai hệ số tải là  $0.640 - 0.523 = 0.117 < 0.2$ . Do vậy, biến DT3 nên được loại bỏ.

Nếu trường hợp biến quan sát cũng tải lên hai hay nhiều nhân tố nhưng chênh lệch hệ số tải lớn hơn 0.2, khi đó biến quan sát nên được giữ lại và xếp biến vào nhân tố có hệ số tải cao hơn.

Cách 1: Loại lần lượt từng biến

- Bước 1: Xác định các biến xấu ở lần phân tích EFA đầu tiên
- Bước 2: Phân loại ba dạng biến xấu 1, 2, 3
- Bước 3: Loại tất cả các biến xấu dạng 1 trước tiên và phân tích lại EFA
  - Bước 4: Nếu tiếp tục xuất hiện biến xấu dạng 1, vẫn loại các biến này trước và chạy tiếp EFA. Nếu không còn biến xấu dạng 1, loại biến xấu dạng 2 và phân tích lại EFA
  - Bước 5: Nếu tiếp tục xuất hiện biến xấu dạng 2, vẫn loại các biến này trước và chạy tiếp EFA. Nếu không còn biến xấu dạng 2, loại biến xấu dạng 3 và phân tích lại EFA.

Cách 2: Loại một lượt các biến xấu trong một lần phân tích EFA

- Bước 1: Xác định các biến xấu ở lần phân tích EFA đầu tiên
- Bước 2: Loại tất cả các dạng biến xấu cùng lúc và phân tích lại EFA
- Bước 3: Nếu tiếp tục xuất hiện biến xấu, loại tất cả các biến đó và tiếp tục phân tích EFA đến khi có kết quả cuối cùng.

Cũng như việc nhận dạng biến xấu để xem xét loại bỏ, hai phương thức loại biến xấu đề xuất ở trên cũng mang tính chất tham khảo để các bạn có thêm căn cứ đưa ra quyết định loại biến như thế nào. Chúng ta có thể linh động việc loại biến một lượt

hay lần lượt, để cuối cùng có được một ma trận xoay tốt nhất và số biến quan sát bị loại bỏ càng ít càng tốt.

**Nếu biến bị loại quá nhiều thì có thể : sử dụng identify outliers để loại những quan sát dị biệt ( bản khảo sát không đi theo xu hướng chung của dữ liệu)**

**Chú ý\*:**

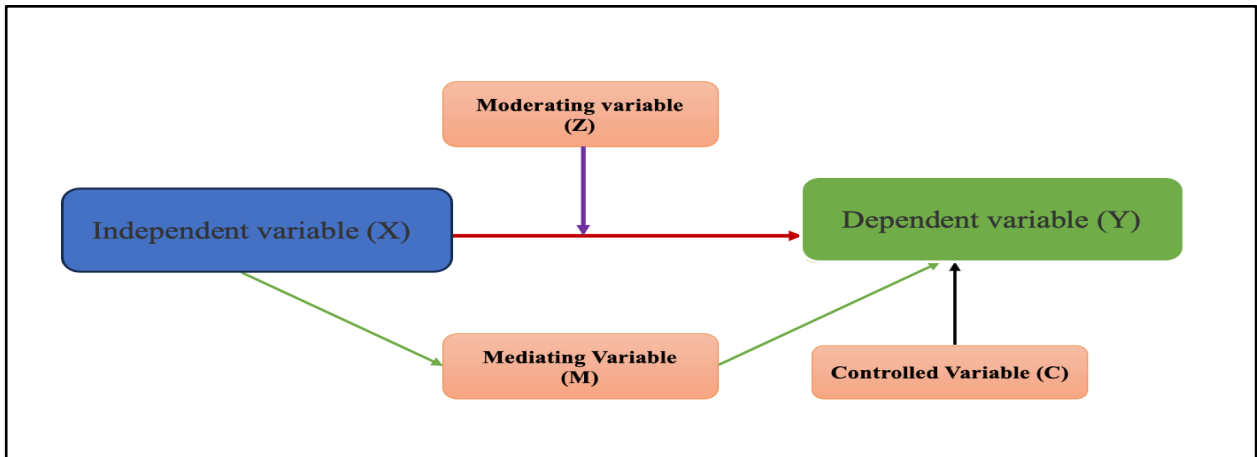
Varimax: xoay vuông góc, sau khi quay trục các nhân tố vẫn ở vị trí vuông góc với nhau. Phép quay này giả định rằng các nhân tố không có sự tương quan với nhau. Phép quay vuông góc ứng dụng nhiều ở các đề tài chỉ có hai loại biến độc lập và phụ thuộc

Promax: xoay không vuông góc, sau khi quay trục các nhân tố sẽ di chuyển đến vị trí phù hợp nhất. Phép quay này giả định các nhân tố có sự tương quan với nhau. Phép quay không vuông góc ứng dụng nhiều ở các đề tài có sự xuất hiện của biến trung gian, lúc này sẽ có các biến vừa đóng vai trò độc lập vừa đóng vai trò phụ thuộc.

Nếu sau EFA, chúng ta đi đến phân phân tích nhân tố khẳng định trên các phần mềm SEM, thì việc sử dụng phép quay Promax cùng phép trích PAF sẽ phù hợp hơn.

Nếu sau EFA, chúng ta đi đến các phân tích tương quan, hồi quy tuyến tính, phép quay Varimax cùng phép trích PCA là một lựa chọn tốt.

**Phân tích EFA cho mô hình có biến trung gian, biến điều tiết, Biến kiểm soát**



Quan điểm của Hair và cộng sự (2010), Quan điểm của Hair và cộng sự (2015)

"Mixing dependent and independent variables in a single factor analysis and then using the derived factors to support dependence relationships is **inappropriate** (**không phù hợp**)"

Khi chạy EFA cho mô hình có biến trung gian, chúng ta sẽ chạy 3 lần EFA cho từng dạng biến: 1 EFA cho độc lập, 1 EFA cho trung gian, 1 EFA cho phụ thuộc.

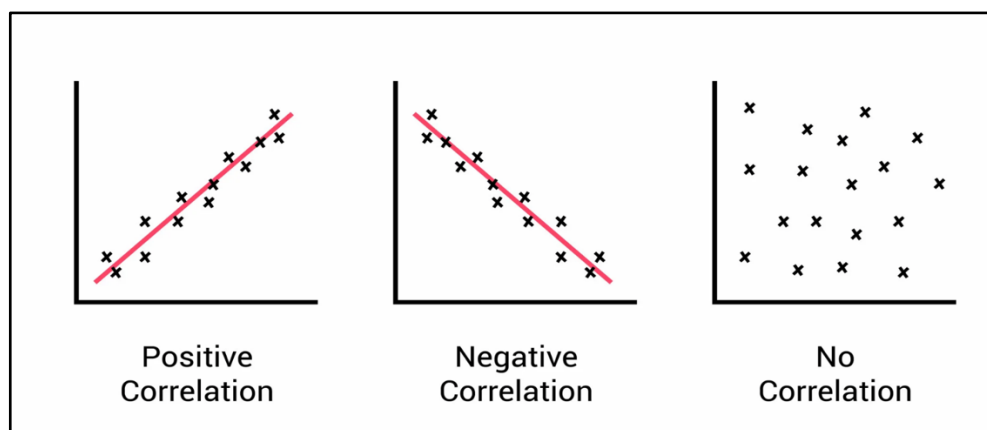
Nếu tồn tại biến điều tiết trong mô hình, chúng ta thực hiện EFA cho riêng một mình biến điều tiết. Nếu có nhiều biến điều tiết cùng điều tiết một mối quan hệ, chúng ta chạy một EFA cho mỗi biến điều tiết. Nếu có nhiều biến điều tiết và mỗi biến điều tiết điều tiết một mối quan hệ khác nhau, chúng ta có thể chạy chung EFA cho tất cả các biến điều tiết.

**Sau khi kết thúc phân tích EFA**, chúng ta sẽ cần chuyển hướng việc đo lường biến quan sát về đo lường nhân tố để kết luận được các giả thuyết đặt ra, bước chuyển đổi này được gọi là tạo nhân tố đại diện.

- Tạo nhân tố đại diện bằng trung bình cộng (mean)
- Tạo nhân tố đại diện bằng tổng giá trị (sum)
- Tạo nhân tố đại diện bằng điểm nhân tố.

#### 1.4. Phân tích tương quan

Tương quan tuyến tính giữa hai biến là mối tương quan mà khi biểu diễn giá trị quan sát của hai biến trên mặt phẳng Oxy, các điểm dữ liệu có xu hướng tạo thành một đường thẳng.



Calculation: <https://www.scribbr.com/statistics/pearson-correlation-coefficient/>

Hệ số tương quan Pearson  $r$  có giá trị dao động từ -1 đến 1:

Khi đã xác định hai biến có mối tương quan tuyến tính (sig nhỏ hơn 0.05), chúng ta sẽ xét đến độ mạnh/yếu của mối tương quan này thông qua trị tuyệt đối của  $r$ . Theo Andy Field (2009):

- $|r| < 0.1$ : mối tương quan rất yếu
- $|r| < 0.3$ : mối tương quan yếu
- $|r| < 0.5$ : mối tương quan trung bình
- $|r| \geq 0.5$ : mối tương quan mạnh

**Kỳ vọng:** sig tương quan giữa độc lập với phụ thuộc nhỏ hơn 0.05 và hệ số tương quan càng cao càng tốt.

**Kỳ vọng:** (1) sig tương quan giữa các biến độc lập lớn hơn 0.05 hoặc (2) sig nhỏ hơn 0.05 và hệ số tương quan sẽ càng thấp càng tốt (nên dưới 0.5).

### **\*\* ý nghĩa 2 dòng cuối trong kết quả pearson**

Khi sig nhỏ hơn 0.05 thì chỗ hệ số tương quan Pearson chúng ta sẽ thấy ký hiệu \* hoặc \*\*.

Ký hiệu \*\* cho biết rằng cặp biến này có sự tương quan tuyến tính ở mức tin cậy đến 99% (tương ứng mức ý nghĩa 1% = 0.01).

Ký hiệu \* cho biết rằng cặp biến này có sự tương quan tuyến tính ở mức tin cậy đến 95% (tương ứng mức ý nghĩa 5% = 0.05).

### 1.5. Phân tích và đọc kết quả hồi quy tuyến tính bội

Hồi quy tuyến tính là phép hồi quy xem xét mối quan hệ tuyến tính – dạng quan hệ đường thẳng giữa biến độc lập với biến phụ thuộc.

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_k X_k + u$$

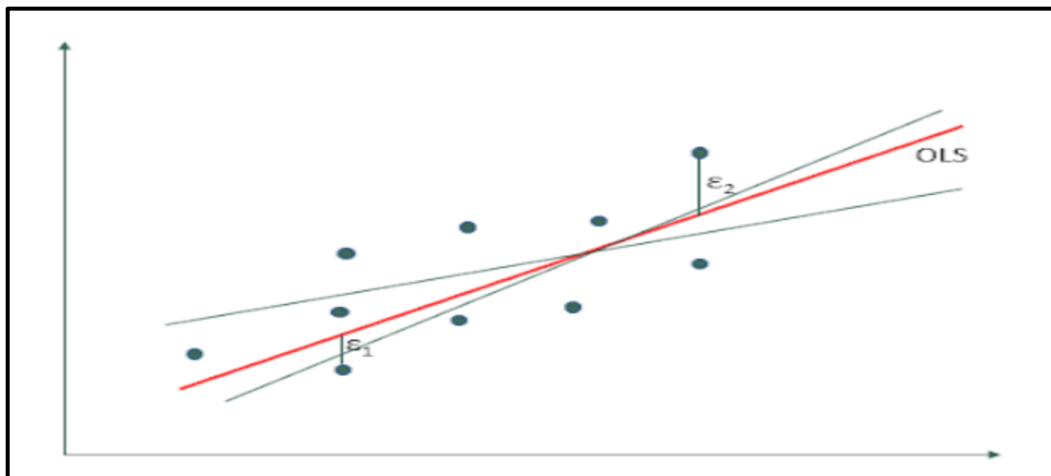
Trong đó:

Y: biến phụ thuộc

$X_1, X_2, X_k$ : biến độc lập

u: Độ lệch chuẩn - Biểu thị sự khác biệt giữa giá trị dự đoán và giá trị thực tế của biến mục tiêu.

$\beta_0, \beta_1, \beta_k$ : là các hệ số hồi quy (hay còn gọi là trọng số), biểu thị mức độ ảnh hưởng của mỗi biến đầu vào đến biến mục tiêu



Nguyên tắc của phương pháp hồi quy OLS là làm cho biến thiên phần dư này trong phép hồi quy là nhỏ nhất. Khi biểu diễn trên mặt phẳng Oxy, đường hồi quy OLS là một đường thẳng đi qua đám đông các điểm dữ liệu mà ở đó, khoảng cách từ các điểm dữ liệu (trị tuyệt đối của  $\varepsilon$ ) đến đường hồi quy là ngắn nhất.

### 1.5.1. Nhận biết phân phối chuẩn

(1) Xem biểu đồ với đường cong chuẩn (Histograms with normal curve) với dạng hình chuông đối xứng với tần số cao nhất nằm ngay giữa và các tần số thấp nằm ở 2 bên. Giá trị trung bình (mean) và trung vị (median) gần bằng nhau và độ nghiêng (Skewness) gần bằng Zero.

(2) vẽ biểu đồ xác suất chuẩn (normal Q- Q Plot). Phân phối chuẩn khi biểu đồ xác suất này có quan hệ tuyến tính (Đường thẳng).

(3) Dùng kiểm định Kolmogorov – smirnov khi cỡ mẫu lớn hơn 50 hoặc phép kiểm Shapiro – wilk khi cỡ mẫu nhỏ hơn 50. Được coi là có phân phối chuẩn khi mức ý nghĩa (p) lớn hơn 0.05

(4) Trong khi thử nghiệm Shapiro – wilk, Kolmogorov – smirnov có thể được sử dụng với mẫu nhỏ hơn 30, và chúng không đáng tin cậy với mẫu lớn hơn. Thử nghiệm độ lệch (Skewness) và độ nhọn kurtosis có thể được sử dụng để xác định độ phân phối chuẩn. Giá trị tuyệt đối của Skewness nhỏ hơn 2 và giá trị tuyệt đối của Kurtosis (proper) nhỏ hơn 7 cho thấy dữ liệu chắc chắn được phân phối chuẩn.

### 1.5.2. Kiểm định đa cộng tuyến

**Hệ số phóng đại phương sai (VIF)** được sử dụng để kiểm tra sự hiện diện của đa cộng tuyến trong mô hình hồi quy.

For a regression model

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n$$

where,

$$VIF = \frac{1}{1 - R^2}$$
$$R^2 = \frac{\sum (y_{\text{calculated}} - \bar{y})^2}{\sum (y_{\text{given}} - \bar{y})^2}$$

Measure of Multicollinearity

1 => không tương quan. Đa cộng tuyến không tồn tại.

Giữa 1 và 5 => tương quan vừa phải. Tồn tại đa cộng tuyến thấp.

Lớn hơn 5 => Tương quan cao. Tồn tại đa cộng tuyến cao.

**Giá trị dung sai tolerance**

Khi  $R^2 = 0$  nghĩa là không có cộng tuyến thì chúng ta có thể nói rằng Dung sai cao (=1).

The inverse of VIF is called **Tolerance** and is given as –

$$TOL = \frac{1}{VIF} = (1 - R^2)$$

Ví dụ: nếu các biến độc lập khác giải thích 25% của biến độc lập X1, thì  $1 - R^2 = 1 - 0.25 = 0.75$ , vậy giá trị dung sai của X1 là 0.75. Giá trị dung sai phải cao, có nghĩa là một mức độ đa cộng tuyến nhỏ (tức là, các biến độc lập khác không có chung một lượng đáng kể nào của phương sai được chia sẻ). Việc xác định các mức dung sai thích hợp sẽ được đề cập trong phần sau.

### 1.5.3. Hiện tượng phương sai thay đổi (Heteroscedasticity)

Phương sai thay đổi (heteroscedasticity) là tình huống thống kê trong đó có sự thay đổi theo một quy luật nào đó trong phần dư hoặc sai số sau khi phương trình hồi quy được ước lượng từ kết quả quan sát mẫu của biến độc lập và phụ thuộc.

Để đánh giá mô hình hồi quy có vi phạm giả định này hay không, chúng ta sẽ dựa vào biểu đồ Scatter Plot giữa các phần dư chuẩn hóa và giá trị dự đoán chuẩn hóa của phép hồi quy. Nếu các điểm dữ liệu tạo thành dạng đám mây có độ lớn trải đều nhau dọc theo đường tung độ 0, chúng ta có thể kết luận giả định phương sai thay đổi không bị vi phạm (thường vùng phân tán của các điểm dữ liệu trên trục tung nằm trong đoạn -3 đến 3 là tốt).

Bên cạnh việc đánh giá hiện tượng phương sai thay đổi qua đồ thị scatter, chúng ta có thể sử dụng các kiểm định như **White, the Breusch-Pagan test, Glesjer, tương quan hạng Spearman** để đánh giá chính xác hơn giả định này.

**Thực hiện White Test để kiểm tra phương sai thay đổi**

**Null ( $H_0$ ):** Homoscedasticity is present.

**Alternative ( $H_A$ ):** Heteroscedasticity is present.

Vì giá trị p không nhỏ hơn 0,05 nên chúng tôi không bác bỏ giả thuyết null. Chúng tôi không có đủ bằng chứng để nói rằng phương sai thay đổi có mặt trong mô hình hồi quy.

Tuy nhiên, nếu bạn bác bỏ giả thuyết null, điều này có nghĩa là phương sai thay đổi có mặt trong dữ liệu.

Có nhiều cách để khắc phục sai phạm phương sai sai số thay đổi, thì phương pháp ước lượng vững ma trận hiệp phương sai (robust standard errors) là được sử dụng nhiều nhất vì tính đơn giản và hiệu quả của nó;

Cách 1: You can try performing a transformation on the response variable, such as taking the log, square root, or cube root of the response variable. Typically this can cause heteroscedasticity to go away.

Cách 2: <https://www.statology.org/weighted-least-squares-in-r/>

### Breusch–Pagan Test

Có thể tham khảo: <https://www.statology.org/breusch-pagan-test-r/>.

#### 1.5.4. Kết quả phân tích hồi quy đa tuyến

- *F-Test*

##### Các giả thuyết cho *F* Test

$$\begin{aligned}H_0 &: \beta_1 = \beta_2 = \dots = \beta_k = 0 \\H_1 &: \text{At least one of } \beta_1, \beta_2, \dots, \beta_k \neq 0\end{aligned}$$

Kiểm tra F bằng cách tìm kiếm giá trị p thích hợp trong việc diễn giải mức ý nghĩa kết quả. Nếu giá trị p lớn hơn 0,05 thì kết quả không có ý nghĩa. Nếu giá trị p nhỏ hơn 0,05, giả thuyết không bị bác bỏ và chúng tôi kết luận rằng mô hình hồi quy có ý nghĩa thống kê. Nếu  $p < 0,01$  thì nó rất có ý nghĩa, v.v.

Ví dụ: Đối với điểm F là 68,428 với số điểm có ý nghĩa là  $0,000 < 0,05$ . Sau đó, mô hình hồi quy có thể được sử dụng để dự báo sự tham gia của các biến. Điều này

cho thấy rằng đồng thời tất cả các biến độc lập đều có tác động đáng kể đến biến phụ thuộc.

```
> x<- cbind(educ,exper,tenure)
> View(x)
> sahd<- lm(lwage~x,data = wage2)
> anova(sahd)
Analysis of Variance Table

Response: lwage
      Df Sum Sq Mean Sq F value    Pr(>F)
x       3  25.695   8.5651  56.974 < 2.2e-16 ***
Residuals 931 139.961   0.1503
---
Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
```

- **Hệ số xác định ( $R^2$ )**

**Hệ số  $R^2$  đã được điều chỉnh (adjusted  $R$  square)**

$$R\text{-squared} = R^2 = \text{SSE}/\text{SST} = 1 - \text{SSR}/\text{SST}$$

Total sum of squares (SST) = explained sum of squares (SSE) + residual sum of squares (SSR)

Công thức tính hệ số  $R^2$  đã được điều chỉnh như sau:

$$\text{Adjusted } R^2 = 1 - [(1 - R^2) * (n - 1) / (n - k - 1)]$$

Trong đó:

- $R^2$  là hệ số  $R^2$
- $n$  là số lượng mẫu trong tập dữ liệu
- $k$  là số lượng biến độc lập trong mô hình.

**Ví dụ minh họa:** Mô hình kinh tế lượng mô tả các yếu tố quyết định tiền lương của người lao động.

$$\widehat{\text{Wage}} = \hat{\beta}_0 + \hat{\beta}_1 \text{education} + \hat{\beta}_2 \text{experience} + \hat{\beta}_3 \text{tenure}$$

Kết quả phân tích hồi quy tuyến tính bội in R

| Source   | SS         | df  | MS         | Number of obs | = | 526    |
|----------|------------|-----|------------|---------------|---|--------|
| Model    | 2194.11162 | 3   | 731.370541 | F(3, 522)     | = | 76.87  |
| Residual | 4966.30269 | 522 | 9.51398982 | Prob > F      | = | 0.0000 |
|          |            |     |            | R-squared     | = | 0.3064 |
|          |            |     |            | Adj R-squared | = | 0.3024 |
| Total    | 7160.41431 | 525 | 13.6388844 | Root MSE      | = | 3.0845 |

| wage   | Coef.     | Std. Err. | t     | P> t  | [95% Conf. Interval] |
|--------|-----------|-----------|-------|-------|----------------------|
| educ   | .5989651  | .0512835  | 11.68 | 0.000 | .4982176 .6997126    |
| exper  | .0223395  | .0120568  | 1.85  | 0.064 | -.0013464 .0460254   |
| tenure | .1692687  | .0216446  | 7.82  | 0.000 | .1267474 .2117899    |
| _cons  | -2.872735 | .7289643  | -3.94 | 0.000 | -4.304799 -1.440671  |

R-squared = SS Model / SS Total = 2194 / 7160 = 1 - SSR/SST = 1 - 4966/7160 = 0.306  
Adj R-squared = 1 - [SSR/(n-k-1)] / [SST/(n-1)] = 1 - (4966/522) / (7160/525) = 0.302  
30% of the variation in wage is explained by the regression and the rest is due to error.

Như vậy, hệ số **adjusted  $R^2$**  bằng 0.3 nghĩa là mô hình hồi qui tuyến tính bội đã xây dựng phù hợp với tập dữ liệu là 30%. Nói cách khác, khoảng 30 % sự thay đổi tiền lương quan sát có thể được giải thích bởi sự khác biệt của 03 thành phần: giáo dục, kinh nghiệm và thâm niên.

Trong nghiên cứu lặp lại, chúng ta thường chọn mức trung gian là 0.5 để phân ra 2 nhánh ý nghĩa mạnh/ý nghĩa yếu và kỳ vọng từ **0.5 đến 1 thì mô hình là tốt, bé hơn 0.5 là mô hình chưa tốt**. when the  $R^2$  value is high, it suggests a better fit for the model.

**Giá trị Durbin–Watson** để đánh giá hiện tượng tự tương quan chuỗi bậc nhất. Giá trị DW = 1.849, nằm trong khoảng 1.5 đến 2.5 nên kết quả không vi phạm giả định tự tương quan chuỗi bậc nhất (Yahua Qiao, 2011)

### • Bảng Coefficients

- + Sig < 0.05: Bác bỏ giả thuyết  $H_0$ , nghĩa là hệ số hồi quy của biến  $X_i$  khác 0 một cách có ý nghĩa thống kê, biến  $X_1$  có tác động lên biến phụ thuộc.
- + Sig > 0.05: Chấp nhận giả thuyết  $H_0$ , nghĩa là hệ số hồi quy của biến  $X_i$  bằng 0 một cách có ý nghĩa thống kê, biến  $X_i$  không tác động lên biến phụ thuộc.

## 1.6. Kiểm định sự khác biệt về tổng thể

Phép kiểm định này được sử dụng trong trường hợp chúng ta cần so sánh trị trung bình về một tiêu chí nghiên cứu nào đó giữa hai đối tượng mà chúng ta quan tâm.

Trước khi thực hiện kiểm định trung bình, ta cần thực hiện kiểm định sự bằng nhau của hai phương sai tổng thể. *Kiểm định này có tên là Levene (Levene's test)*, với giả thiết  $H_0$  rằng phương sai của hai tổng thể bằng nhau.

Với mức ý nghĩa 0,05, chúng ta có:

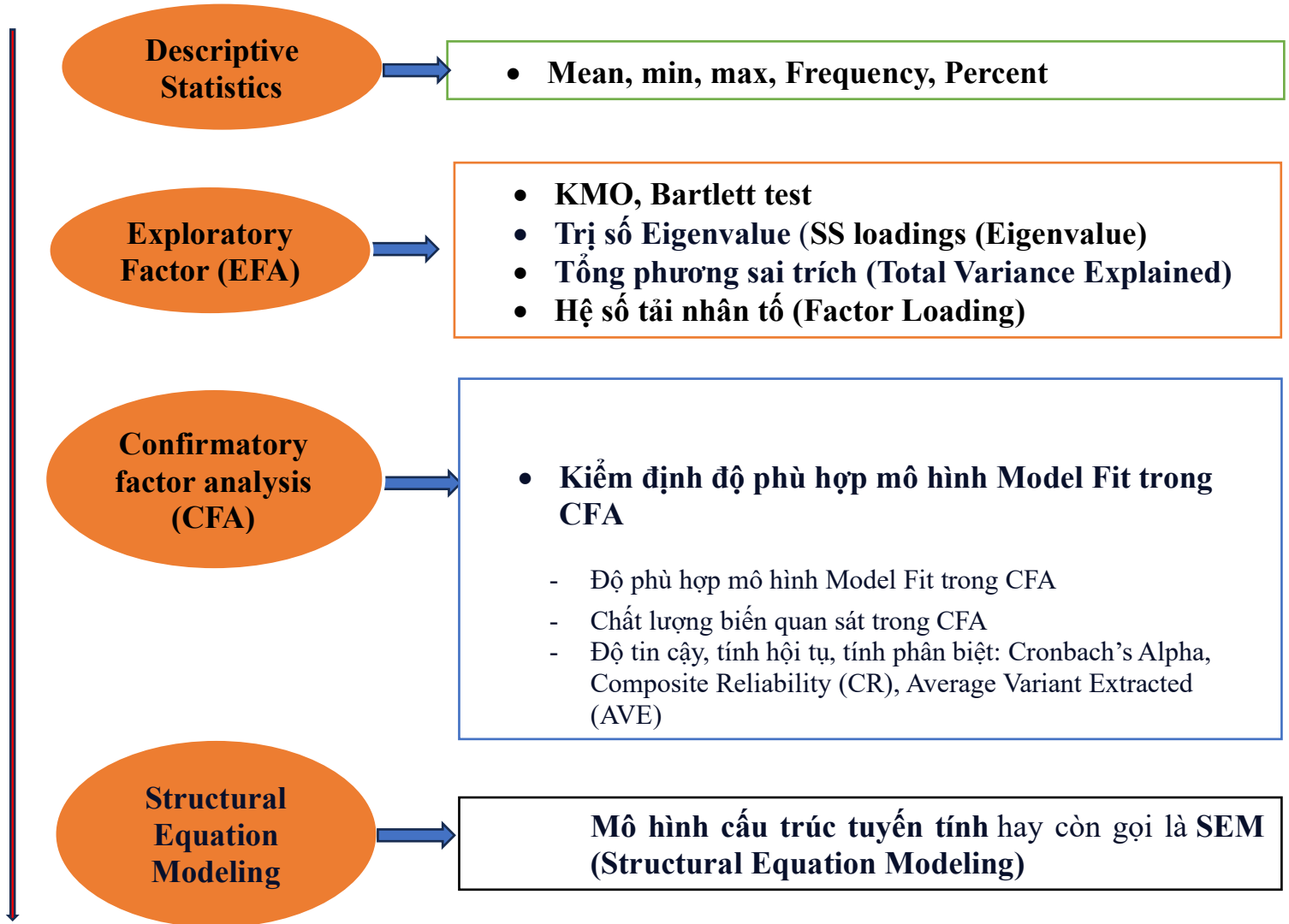
- Nếu  $\text{Sig} < 0.05$ : phương sai giữa các nhóm không có sự khác biệt.
- Nếu  $\text{Sig} \geq 0.05$ : phương sai giữa các nhóm có sự khác biệt.

Kết quả của việc chấp nhận hay bác bỏ  $H_0$  ảnh hưởng quan trọng đến việc chúng ta sẽ lựa chọn loại kiểm định giả thiết về sự bằng nhau của hai trung bình tổng thể nào: kiểm định trung bình với phương sai bằng nhau hay kiểm định trung bình với phương sai khác nhau.

Trong trường hợp  $p\text{-value} < 0.05$  thì giả định  $H_0$  được chấp nhận điều đó có nghĩa là phương sai giữa các nhóm có sự khác biệt. Do đó kết quả ở bảng ANOVA không sử dụng được vì thế phải sử dụng kiểm định Kruskal-Wallis. Còn  $p\text{-value} > 0.05$  thì phương sai giữa các nhóm không sự khác biệt và kết quả ở bảng ANOVA sẽ được sử dụng cho các giá trị này.

Kiểm định Kruskal-Wallis là một kiểm định phi tham số được sử dụng để kiểm tra sự khác biệt về giá trị trung vị giữa các nhóm dữ liệu độc lập. Nếu  $p\text{-value}$  được tính toán bởi kiểm định Kruskal-Wallis là nhỏ hơn một ngưỡng xác định trước (thường là 0,05), thì chúng ta có thể kết luận rằng có sự khác biệt đáng kể giữa các nhóm dữ liệu.

## 2. HƯỚNG DẪN PHƯƠNG PHÁP NGHIÊN CỨU ĐỊNH LƯỢNG CHO DỰ ÁN KHOA HỌC ĐỐI VỚI R VÀ SMART PLS 4



### 2.1: Thống kê mô tả (Descriptive Statistics)

Thống kê mô tả là các hệ số thông tin ngắn gọn tóm tắt một tập dữ liệu nhất định, có thể là đại diện cho toàn bộ tổng thể hoặc một mẫu của tổng thể.

| Lí thuyết                                                        | Hàm R     |
|------------------------------------------------------------------|-----------|
| Số trung bình: $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$          | mean (x)  |
| Phương sai: $s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$ | var (x)   |
| Độ lệch chuẩn: $s = \sqrt{s^2}$                                  | sd (x)    |
| Sai số chuẩn (standard error): $SE = \frac{s}{\sqrt{n}}$         | Không có  |
| Trị số thấp nhất                                                 | min (x)   |
| Trị số cao nhất                                                  | max (x)   |
| Toàn cự (range)                                                  | range (x) |

### Chú ý :

Độ lệch chuẩn càng cao thì hiệu thị sự chênh lệch càng nhiều trong đáp áp của người khảo sát.

### Interquartile Range

- The upper fence = Q3 + IQR\*1.5
- The lower fence = Q1 - IQR\*1.5
- IQR = Q3 – Q1

Thống kê các biến kiểm soát trong nghiên cứu với Frequency, Percent

## 2.2. Phân tích nhân tố khám phá (EFA)

**Phân tích nhân tố khám phá**, gọi tắt là EFA, dùng để rút gọn một tập hợp k biến quan sát thành một tập F (với  $F < k$ ) các nhân tố có ý nghĩa hơn. Với kiểm định độ tin cậy thang đo Cronbach Alpha, chúng ta đang đánh giá mối quan hệ giữa các biến trong cùng một nhóm, cùng một nhân tố, chứ không xem xét mối quan hệ giữa tất cả các biến quan sát ở các nhân tố khác. Trong khi đó, EFA xem xét mối quan hệ giữa các biến ở tất cả các nhóm (các nhân tố) khác nhau nhằm phát hiện ra những biến quan sát tải lên nhiều nhân tố hoặc các biến quan sát bị phân sai nhân tố từ ban đầu.

### Các tiêu chí trong phân tích EFA

- **Hệ số KMO (Kaiser-Meyer-Olkin)** là một chỉ số dùng để xem xét sự thích hợp của phân tích nhân tố. Trị số của KMO phải đạt giá trị 0.5 trở lên ( $0.5 \leq KMO \leq 1$ ) là điều kiện đủ để phân tích nhân tố là phù hợp. Nếu trị số này nhỏ hơn 0.5, thì phân tích nhân tố có khả năng không thích hợp với tập dữ liệu nghiên cứu.
- **Kiểm định Bartlett (Bartlett's test of sphericity)** dùng để xem xét các biến quan sát trong nhân tố có tương quan với nhau hay không. Kiểm định Bartlett có ý nghĩa thống kê ( $\text{sig Bartlett's Test} < 0.05$ ), chứng tỏ các biến quan sát có tương quan với nhau trong nhân tố.
- **Trị số Eigenvalue** là một tiêu chí sử dụng phổ biến để xác định số lượng nhân tố trong phân tích EFA. Với tiêu chí này, chỉ có những nhân tố nào có  $\text{Eigenvalue} > 1$  mới được giữ lại trong mô hình phân tích.
- **Tổng phương sai trích (Total Variance Explained)**  $\geq 50\%$  cho thấy mô hình EFA là phù hợp. Coi biến thiên là 100% thì trị số này thể hiện các nhân tố được trích cô đọng được bao nhiêu % và bị thất thoát bao nhiêu % của các biến quan sát.
- **Hệ số tải nhân tố (Factor Loading)** hay còn gọi là trọng số nhân tố, giá trị này biểu thị mối quan hệ tương quan giữa biến quan sát với nhân tố.

Theo **Hair và cộng sự (2010), Multivariate Data Analysis** hệ số tải từ 0.5 là biến quan sát đạt chất lượng tốt, tối thiểu nên là 0.3.

Hair và cộng sự cũng cho rằng, giá trị tiêu chuẩn của hệ số tải Factor Loading nên được xem xét cùng kích thước mẫu.

| <b>Giá trị Factor Loading</b> | <b>Kích thước mẫu tối thiểu có ý nghĩa thống kê</b> |
|-------------------------------|-----------------------------------------------------|
| <b>0.3</b>                    | <b>350</b>                                          |
| <b>0.35</b>                   | <b>250</b>                                          |
| <b>0.40</b>                   | <b>200</b>                                          |
| <b>0.45</b>                   | <b>150</b>                                          |
| <b>0.50</b>                   | <b>120</b>                                          |
| <b>0.55</b>                   | <b>100</b>                                          |
| <b>0.60</b>                   | <b>85</b>                                           |
| <b>0.65</b>                   | <b>70</b>                                           |
| <b>0.70</b>                   | <b>60</b>                                           |
| <b>0.75</b>                   | <b>50</b>                                           |

Các dạng biến xấu trong EFA

| <b>Rotated component matrix<sup>a</sup></b> |             |             |             |          |          |
|---------------------------------------------|-------------|-------------|-------------|----------|----------|
| <b>Component</b>                            |             |             |             |          |          |
|                                             | <b>1</b>    | <b>2</b>    | <b>3</b>    | <b>4</b> | <b>5</b> |
| <b>LD1</b>                                  | <b>.825</b> |             |             |          |          |
| <b>LD2</b>                                  | <b>.800</b> |             |             |          |          |
| <b>LD3</b>                                  | <b>.755</b> |             |             |          |          |
| <b>DT3</b>                                  | <b>.640</b> | <b>.523</b> |             |          |          |
| <b>DT1</b>                                  |             | <b>.822</b> |             |          |          |
| <b>DT4</b>                                  |             | <b>.796</b> |             |          |          |
| <b>DT2</b>                                  |             | <b>.748</b> |             |          |          |
| <b>DN1</b>                                  |             |             | <b>.836</b> |          |          |

|            |  |  |             |             |             |
|------------|--|--|-------------|-------------|-------------|
| <b>DN3</b> |  |  | <b>.825</b> |             |             |
| <b>DN2</b> |  |  | <b>.796</b> |             |             |
| <b>TL4</b> |  |  |             | <b>.769</b> |             |
| <b>TL3</b> |  |  |             | <b>.760</b> |             |
| <b>TL1</b> |  |  |             | <b>.749</b> |             |
| <b>DN4</b> |  |  |             |             | <b>.899</b> |
| <b>TL2</b> |  |  |             | <b>.421</b> |             |

### **Trường hợp 1: Biến quan sát không đảm bảo hệ số tải tiêu chuẩn**

Tác giả chọn ngưỡng hệ số tải là 0.5 và quy định cho ma trận xoay chỉ hiển thị các biến quan sát có hệ số tải từ 0.4 trở lên (các hệ số tải nhỏ hơn 0.4 sẽ bị ẩn đi). Biến TL2 trong bảng Rotated Component Matrix có hệ số tải cao nhất là  $0.421 < 0.5$  (hệ số tải biến TL2 ở các ô còn lại dưới 0.4 nên đã bị ẩn đi), biến này sẽ được loại bỏ.

### **Trường hợp 2: Biến quan sát chỉ nằm tách biệt một mình ở một nhân tố**

Trong bảng ma trận xoay ở ví dụ trên, biến DN4 nằm tách biệt một mình ở nhân tố số 5. Chúng ta nên loại bỏ biến quan sát này đi.

### **Trường hợp 3: Tải lên nhiều nhóm nhân tố và chênh lệch hệ số tải nhỏ hơn 0.2**

Matt C. Howard (2015) cho rằng, nếu một biến quan sát tải lên ở hai nhân tố nhưng chênh lệch hệ số tải nhỏ hơn 0.2, biến quan sát nên được xem xét loại bỏ.

Trong bảng ma trận xoay ở ví dụ trên, biến DT3 tải lên hai nhân tố số 1 và số 2 với hệ số tải lần lượt ở hai nhân tố là 0.640 và 0.523. Hiệu số hai hệ số tải là  $0.640 - 0.523 = 0.117 < 0.2$ . Do vậy, biến DT3 nên được loại bỏ.

Nếu trường hợp biến quan sát cũng tải lên hai hay nhiều nhân tố nhưng chênh lệch hệ số tải lớn hơn 0.2, khi đó biến quan sát nên được giữ lại và xếp biến vào nhân tố có hệ số tải cao hơn.

**Chú ý\*:**

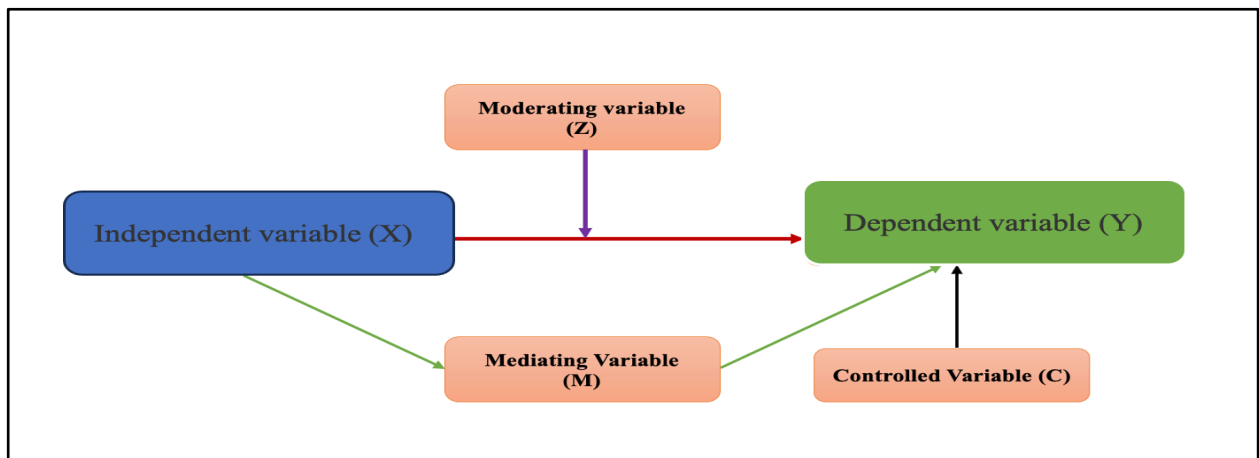
Varimax: xoay vuông góc, sau khi quay trục các nhân tố vẫn ở vị trí vuông góc với nhau. Phép quay này giả định rằng các nhân tố không có sự tương quan với nhau. Phép quay vuông góc ứng dụng nhiều ở các đề tài chỉ có hai loại biến độc lập và phụ thuộc

Promax: xoay không vuông góc, sau khi quay trục các nhân tố sẽ di chuyển đến vị trí phù hợp nhất. Phép quay này giả định các nhân tố có sự tương quan với nhau. Phép quay không vuông góc ứng dụng nhiều ở các đề tài có sự xuất hiện của biến trung gian, lúc này sẽ có các biến vừa đóng vai trò độc lập vừa đóng vai trò phụ thuộc.

**Nếu sau EFA, chúng ta đi đến phân phân tích nhân tố khẳng định trên các phần mềm SEM, thì việc sử dụng phép quay Promax cùng phép trích PAF (Principal Axis Factoring) sẽ phù hợp hơn.**

Nếu sau EFA, chúng ta đi đến các phân tích tương quan, hồi quy tuyến tính, phép quay Varimax cùng phép trích PCA là một lựa chọn tốt.

### **Phân tích EFA cho mô hình có biến trung gian, biến điều tiết, Biến kiểm soát**



Quan điểm của Hair và cộng sự (2010), Quan điểm của Hair và cộng sự (2015)

"Mixing dependent and independent variables in a single factor analysis and then using the derived factors to support dependence relationships is **inappropriate** (**không phù hợp**)"

Khi chạy EFA cho mô hình có biến trung gian, chúng ta sẽ chạy 3 lần EFA cho từng dạng biến: 1 EFA cho độc lập, 1 EFA cho trung gian, 1 EFA cho phụ thuộc.

Nếu tồn tại biến điều tiết trong mô hình, chúng ta thực hiện EFA cho riêng một mình biến điều tiết. Nếu có nhiều biến điều tiết cùng điều tiết một mối quan hệ, chúng ta chạy một EFA cho mỗi biến điều tiết. Nếu có nhiều biến điều tiết và mỗi biến điều tiết điều tiết một mối quan hệ khác nhau, chúng ta có thể chạy chung EFA cho tất cả các biến điều tiết.

### **2.3. Confirmatory factor analysis (CFA)**

Phân tích nhân tố khẳng định (Confirmatory Factor Analysis - CFA) là một loại mô hình cấu trúc tuyến tính (SEM) tập trung vào mô hình đo lường (measurement models), cụ thể là mối quan hệ giữa các biến quan sát hoặc chỉ báo (indicators) với các biến tiềm ẩn (latent variables) hoặc còn gọi là nhân tố (factors).

Thực tế, các giả định này hiếm khi đạt được trong các nghiên cứu thuộc khoa học hành vi và khoa học xã hội, trong các nghiên cứu này các biến số được thu thập dựa trên bảng câu hỏi khảo sát và các quan sát độc lập, hoặc các hình thức tương tự. Do đó, các kết quả thu được từ phương pháp bình phương nhỏ nhất (ví dụ, tương quan Pearson, hồi quy tuyến tính) thường không cho biết mức độ của sai số đo lường ở các biến sử dụng trong phân tích. Trong khi đó, phân tích nhân tố CFA và phân tích mô hình cấu trúc SEM cho phép ước lượng mối quan hệ giữa các biến số trong mô hình sau khi đã điều chỉnh sai số đo lường và sai số lý thuyết, do vậy các kết quả đầu ra sẽ mang tính "thực" cao đáng kể.

#### **2.3.1. Kiểm định độ phù hợp mô hình Model Fit trong CFA**

Theo Hu & Bentler (1999), Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives, Structural Equation Modeling các chỉ số được xem xét để đánh giá Model Fit phổ biến gồm (các nhà nghiên cứu thường sử dụng nguồn này):

- $\text{CMIN/df} \leq 3$  là tốt,  $\text{CMIN/df} \leq 5$  là chấp nhận được
- $\text{CFI} \geq 0.9$  là tốt,  $\text{CFI} \geq 0.95$  là rất tốt,  $\text{CFI} \geq 0.8$  là chấp nhận được (CFA dao động trong vùng 0 đến 1)
- $\text{GFI} \geq 0.9$  là tốt,  $\text{GFI} \geq 0.95$  là rất tốt (chỉ số này bạn cần đọc thêm phần lưu ý cuối bài viết)
- $\text{TLI} \geq 0.9$  là tốt
- $\text{RMSEA} \leq 0.06$  là tốt,  $\text{RMSEA} \leq 0.08$  là chấp nhận được
- $\text{PCLOSE} \geq 0.05$  là tốt,  $\text{PCLOSE} \geq 0.01$  là chấp nhận được

Theo Hair et al. (2010), Multivariate Data Analysis, 7th edition các chỉ số được xem xét để đánh giá Model Fit gồm:

- $\text{CMIN/df} \leq 2$  là tốt,  $\text{CMIN/df} \leq 5$  là chấp nhận được
- $\text{CFI} \geq 0.9$  là tốt,  $\text{CFI} \geq 0.95$  là rất tốt,  $\text{CFI} \geq 0.8$  là chấp nhận được (CFA dao động trong vùng 0 đến 1)
- $\text{GFI} \geq 0.9$  là tốt,  $\text{GFI} \geq 0.95$  là rất tốt
- $\text{RMSEA} \leq 0.08$  là tốt,  $\text{RMSEA} \leq 0.03$  là rất tốt

Trường hợp riêng của GFI lớn hơn 0.8 và nhỏ hơn 0.9, Nếu bạn bắt buộc phải đánh giá chỉ số này trong bài luận, nhưng ngưỡng của nó dưới 0.9, bạn có thể chấp nhận ngưỡng 0.8 theo 2 công trình nghiên cứu của Baumgartner and Homburg (1995) và của Doll, Xia, and Torkzadeh (1994). Nguồn trích dẫn:

Với các trường hợp các chỉ số không đạt được các ngưỡng chấp nhận, để tăng độ phù hợp mô hình Model Fit trong CFA, chúng ta thường dùng mũi tên 2 chiều Covariance trong AMOS để nối các sai số có chỉ số MI (Modification Indices) cao trong cùng một thang đo lại với nhau. Điều này sẽ giúp giảm Chi-square và cải thiện

các chỉ số Model Fit. Tuy nhiên, không nên lạm dụng nói quá nhiều các sai số sẽ dẫn đến những sai lệch về tính chính xác của dữ liệu.

### 2.3.2 Chất lượng biến quan sát trong CFA

Hệ số standardized regression weight còn được sử dụng như một giá trị đánh giá mức độ đóng góp của biến quan sát lên biến tiềm ẩn. Biến quan sát nào có standardized regression weight lớn hơn thì sẽ đóng góp vào biến mẹ nhiều hơn.

Theo Hair và cộng sự (2010), những biến quan sát có **standardized regression weight tối thiểu từ 0.5 trở lên** sẽ được giữ lại, lý tưởng nhất là từ 0.7 trở lên (standardized regression weight ở CFA AMOS trong một số tài liệu còn được gọi là factor loading). Những biến có standardized regression weight nhỏ hơn 0.5 cần được loại bỏ (khi có biến quan sát được loại bỏ, chúng ta sẽ quay lại diagram CFA xóa biến quan sát này khỏi mô hình, thực hiện phân tích lại CFA và đánh giá chất lượng biến quan sát lại một lần nữa).

## 2.4 Độ tin cậy, tính hội tụ, tính phân biệt

### a. Cronbach's Alpha

Cronbach's Alpha là một cách để đo lường tính nhất quán bên trong của một bảng câu hỏi hoặc khảo sát.

Các hệ số

- Cronbach's Alpha: Hệ số Cronbach's Alpha
- N of Items: Số lượng biến quan sát
- Scale Mean if Item Deleted: Trung bình thang đo nếu loại biến
- Scale Variance if Item Deleted: Phương sai thang đo nếu loại biến
- Corrected Item-Total Correlation: Tương quan biến tổng
- Cronbach's Alpha if Item Deleted: Hệ số Cronbach's Alpha nếu loại biến

**Cronbach's Alpha nằm trong khoảng từ 0 đến 1, với các giá trị càng cao cho thấy khảo sát hoặc bảng câu hỏi đáng tin cậy hơn.**

Bảng sau đây mô tả các giá trị khác nhau của Cronbach's Alpha thường được diễn giải như sau:

| Cronbach's Alpha        | Internal consistency |
|-------------------------|----------------------|
| $0.9 \leq \alpha$       | Excellent            |
| $0.8 \leq \alpha < 0.9$ | Good                 |
| $0.7 \leq \alpha < 0.8$ | Acceptable           |
| $0.6 \leq \alpha < 0.7$ | Questionable         |
| $0.5 \leq \alpha < 0.6$ | Poor                 |
| $\alpha < 0.5$          | Unacceptable         |

Theo Nunnally (1978), một thang đo tốt nên có độ tin cậy Cronbach's Alpha từ 0.7 trở lên. Hair và cộng sự (2009) cũng cho rằng, một thang đo đảm bảo tính đơn hướng và đạt độ tin cậy nên đạt ngưỡng Cronbach's Alpha từ 0.7 trở lên, tuy nhiên, với tính chất là một nghiên cứu khám phá sơ bộ, ngưỡng Cronbach's Alpha là 0.6 có thể chấp nhận được. Hệ số Cronbach's Alpha càng cao thể hiện độ tin cậy của thang đo càng cao.

**Corrected Item - Total Correlation** thể hiện mối quan hệ giữa biến quan sát đó với tất cả các biến quan sát còn lại trong cùng 1 thang đo (cùng 1 nhóm).

#### Các ngưỡng cần chú ý:

- 0.30 by Nunnally, J. C., & Bernstein, I. H. (1994)
- 0.30 by Cristobal et al. (2007)
- 0.40 by Loiacono et al. (2002)
- 0.50 by Francis and White (2002) and Kim and Stoel (2004)

#### Chú ý\*

- Nếu thang đo có hệ số Cronbach's Alpha dưới 0.6, chúng ta **chưa vội** kết luận thang đo không có độ tin cậy mà cần kiểm tra và loại hết biến quan sát

có giá trị Cronbach's Alpha if Item Deleted cao hơn mức 0.6. Đến khi loại hết rồi mà hệ số Cronbach's Alpha vẫn dưới 0.6 thì mới kết luận.

- Nếu hệ số Cronbach's Alpha của nhóm đã **đủ tiêu** chuẩn thì việc xuất hiện biến quan sát có Cronbach's Alpha if Item Deleted lớn hơn Cronbach's Alpha của nhóm nhưng tương quan biến tổng lớn hơn 0.3 thì chúng ta không cần loại biến quan sát đó đi.
- Nếu hệ số Cronbach's Alpha của nhóm **chưa đủ tiêu** chuẩn thì việc xuất hiện biến quan sát có Cronbach's Alpha if Item Deleted lớn hơn Cronbach's Alpha của nhóm nhưng tương quan biến tổng lớn hơn 0.3 thì chúng ta nên loại biến quan sát đó đi để cải thiện độ tin cậy thang đo cho tới khi hệ số Cronbach's Alpha của nhóm đạt tiêu chuẩn.
- Nếu hệ số Cronbach's Alpha của nhóm **chưa đủ tiêu chuẩn**, chúng ta đã loại các biến quan sát có Cronbach's Alpha if Item Deleted lớn hơn Cronbach's Alpha của nhóm nhưng thang đo vẫn không đủ tiêu chuẩn. Khi đó, thang đo không đảm bảo độ tin cậy cho nghiên cứu, cần loại bỏ cả thang đo này.
- Nếu sự chênh lệch giữa Cronbach's Alpha của nhóm với Cronbach's Alpha if Item Deleted của biến quan sát là đáng kể từ **0.3 trở lên**. Chúng ta sẽ loại biến quan sát đó để tăng thêm độ tin cậy của thang đo.

Theo Hair và cộng sự (2010) và Hair và cộng sự (2016) chúng ta sử dụng các chỉ số CR, AVE, MSV, bảng Fornell and Larcker để đánh giá tính hội tụ, tính phân biệt thang đo.

#### **b. Tính hội tụ (Convergent Validity)**

- Độ tin cậy tổng hợp Composite Reliability (CR)  $\geq 0.7$ .
- Phương sai trung bình được trích Average Variance Extracted (AVE)  $\geq 0.5$ .

#### **c. Tính phân biệt (Discriminant Validity)**

- Phương sai chia sẻ lớn nhất Maximum Shared Variance (MSV) < Average Variance Extracted (AVE)
- Căn bậc hai phương sai trung bình được trích Square Root of AVE (SQRTAVE) > Tương quan giữa các cấu trúc Inter-Construct Correlations trong bảng Fornell and Larcker.

## 2.4. Mô hình cấu trúc tuyến tính - SEM (Structural Equation Modeling)

**Mô hình cấu trúc tuyến tính** hay còn gọi là **SEM (Structural Equation Modeling)** là một kỹ thuật phân tích thống kê thể hệ thứ hai được phát triển để phân tích mối quan hệ đa chiều giữa nhiều biến trong một mô hình (Haenlein & Kaplan, 2004).

SEM là một kỹ thuật thống kê mạnh mẽ hơn để giải quyết các yêu cầu sau:

1. Phân tích nhiều mô hình hồi quy bởi một cách đồng thời;
2. Phân tích hồi quy với bài toán đa cộng tuyến;
3. Phân tích phân tích đường dẫn với nhiều biến phụ thuộc;
4. Mô hình hóa mối quan hệ đa chiều giữa các biến trong một mô hình.

### **Một số lưu ý:**

#### **1. Biến bậc hai (second-order) và biến bậc một (first-order)**

- Trong thuật ngữ thống kê chúng ta gọi là **biến bậc hai - biến bậc một**. Còn trong trao đổi, nói chuyện chúng ta có thể sử dụng thêm cụm biến mẹ - biến con hoặc biến mẹ - biến thành phần để hiểu một cách trực quan hơn đó là một biến mẹ sẽ chứa các biến con bên trong.
- Biến bậc 2, không có biến quan sát. Lưu ý rằng, biến bậc 2 được đo lường thông qua các biến bậc 1, do đó biến bậc 2 không có biến quan sát.
- Biến bậc 1, gồm các biến quan sát, là biến thành phần của biến Biến bậc

2

#### **1. Theo Bollen (2011), có 3 loại thang đo đo lường gồm:**

- **Effect Indicator (Reflective Measurement): mô hình Reflective**

Đây là dạng đo lường mà các biến quan sát là kết quả được tạo ra từ biến tiềm ẩn.

Ví dụ HL1: Cảm thấy thoải mái khi mỗi ngày đến công ty làm việc

- **Composite Indicator (Formative Measurement): mô hình Formative**

Đây là dạng đo lường mà các biến quan sát là nguyên nhân tạo ra biến tiềm ẩn.

Ví dụ HL1: Tôi hài lòng về điều kiện làm việc ở công ty

**Nội dung cần nắm về Formative:**

1. Biến quan sát là các thành phần cấu thành nên biến tiềm ẩn
2. Các biến quan sát ít có sự tương quan với nhau
3. Khi bỏ 1 biến quan sát khỏi mô hình thang đo, biến tiềm ẩn bị ảnh hưởng rất nhiều
4. Các biến quan sát không thể thay thế vai trò cho nhau
5. **KHÔNG THỂ sử dụng các phân tích: Cronbach Alpha, Composite Reliability, EFA, CFA**

Trên đây chính là điểm khác biệt giữa 2 thang đo Formative và Reflective, các bạn cần chú ý sử dụng một cách hợp lý. Bởi một số bạn xây dựng thang đo biến quan sát cho biến tiềm ẩn dạng Formative nhưng lại thực hiện phân tích độ tin cậy thang đo Cronbach Alpha, phân tích nhân tố khám phá EFA, dẫn đến kết quả độ tin cậy thang đo không đạt ngưỡng chấp nhận, ma trận xoay lộn xộn. Các bạn lại hoang mang nghi ngờ về số liệu thu thập có vấn đề, tuy nhiên, vấn đề lại nằm ở thang đo của bạn không phù hợp để thực hiện các kiểm định này.

- **Causal Indicator (Causal Measurement): mô hình Causal**

**3. Bài nghiên cứu sử dụng Cronbach Alpha, EFA, CFA, SEM. Trong đó:**

- Cronbach Alpha yêu cầu mẫu tối thiểu là 100
- EFA yêu cầu mẫu tối thiểu là 200
- SEM yêu cầu mẫu tối thiểu là 150

Chúng ta sẽ lấy mẫu tối thiểu cho bài nghiên cứu dựa theo EFA là 200.

Đối với phân tích mô hình cấu trúc tuyến tính SEM, theo **Hair et al., *Multivariate Data Analysis*, 2010, 7th edition, p.574** thì cỡ mẫu phù hợp sẽ được xác định dựa trên các nhóm nhân tố. Cụ thể:

- **Cỡ mẫu tối thiểu là 100** - Số cấu trúc tiềm ẩn từ **5 nhóm trở xuống**, mỗi nhóm phải có hơn 3 biến quan sát.
- **Cỡ mẫu tối thiểu là 150** - Số nhóm nhân tố từ **7 nhóm trở xuống**, mỗi nhóm phải có từ 3 biến quan sát, communality của các biến quan sát phải từ 0.5 trở lên.
- **Cỡ mẫu tối thiểu là 300** - Số nhóm nhân tố từ **7 nhóm trở xuống**, mỗi nhóm phải có từ 3 biến quan sát, communality của các biến quan sát có thể chỉ cần từ 0.45 trở lên.
- **Cỡ mẫu tối thiểu là 500** - Số nhóm nhân tố trên 7, mỗi nhóm có thể có ít hơn 3 biến quan sát.

## TÀI LIỆU THAM KHẢO

- 1) Wooldridge, Jeffrey M., 1960-. (2012). Introductory econometrics: a modern approach. Mason, Ohio: South-Western Cengage Learning,
- 2) Christoph Hanck.ect (2020). Introduction to Econometrics with R, Department of Business Administration and Economics University of Duisburg-Essen Essen, Germany
- 3) phamlocblog “các bài báo về phân tích hồi quy bội,” Thời gian tài liệu được tạo hay cập nhật. [Trực tuyến]. Địa chỉ: <https://www.phamlocblog.com/p/dich-vu-spss.html> [Truy cập 20/10/2023]

**END**

